

EDSA-Ensemble: An Event Detection Sentiment Analysis Ensemble Architecture

Alexandru Petrescu^{ID}, Ciprian-Octavian Truică^{ID}, Elena-Simona Apostol^{ID}, and Adrian Paschke^{ID}

Abstract—As global digitization continues to grow, technology becomes more affordable and easier to use, and social media platforms thrive, becoming the new means of spreading information and news. Communities are built around sharing and discussing current events. Within these communities, users are enabled to share their opinions about each event. Using Sentiment Analysis to understand the polarity of each message belonging to an event, as well as the entire event, can help to better understand the general and individual feelings of significant trends and the dynamics on online social networks. In this context, we propose a new ensemble architecture, EDSA-Ensemble (Event Detection Sentiment Analysis Ensemble), that uses Event Detection and Sentiment Analysis to improve the detection of the polarity for current events from Social Media. For Event Detection, we use techniques based on Information Diffusion taking into account both the time span and the topics. To detect the polarity of each event, we preprocess the text and employ several Machine and Deep Learning models to create an ensemble model. The preprocessing step includes several word representation models: raw frequency, *TFIDF*, Word2Vec, and Transformers. The proposed EDSA-Ensemble architecture improves the event sentiment classification over the individual Machine and Deep Learning models.

Index Terms—Ensemble model, event detection, sentiment analysis, social networks analysis, event sentiment analysis.

Manuscript received 14 June 2023; revised 22 July 2024; accepted 23 July 2024. Date of publication 9 August 2024; date of current version 27 May 2025. This work was supported in part by the National Program for Research of the National Association of Technical Universities under Grant GNAC ARUT 2023 through the project “StopXocial: Smart system for analyzing and stopping antisocial behavior in online communities” (Contract no. 64/10.10.2023), in part by the German Academic Exchange Service (DAAD) through the project “iTracing: Automatic Misinformation Fact-Checking” under DAAD Grant 91809005, and in part by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) Project RECOMP (DFG – GZ: PA 1820/5-1) and the Institute of Applied Informatics (InfAI). Recommended for acceptance by J. Epps. (Alexandru Petrescu, Ciprian-Octavian Truică, and Elena-Simona Apostol are contributed equally to this work.) (Corresponding author: Ciprian-Octavian Truică.)

Alexandru Petrescu and Ciprian-Octavian Truică are with the Computer Science and Engineering Department, Faculty of Automatic Control and Computers, National University of Science and Technology Politehnica Bucharest, 060042 Bucharest, Romania (e-mail: ciprian.truica@upb.ro).

Elena-Simona Apostol is with the Computer Science and Engineering Department, Faculty of Automatic Control and Computers, National University of Science and Technology Politehnica Bucharest, 060042 Bucharest, Romania, and also with Academy of Romanian Scientists, 3 Ilfov Bucharest, Romania (e-mail: elena.apostol@upb.ro).

Adrian Paschke is with the Fraunhofer Institute for Open Communication Systems, Freie Universitaet Berlin, 10589 Berlin, Germany, also with the Freie Universitaet Berlin, 14195 Berlin, Germany, and also with the Institute of Applied Informatics (InfAI), 04109 Leipzig, Germany (e-mail: adrian.paschke@fokus.fraunhofer.de).

Digital Object Identifier 10.1109/TAFFC.2024.3434355

I. INTRODUCTION

AS SOCIAL media platforms grow more and more each day, it also increases the need to analyze and understand certain aspects, such as the impact of important or spiking topics over the network [51]. Event Detection techniques are used to automatically identify important or spiking topics (i.e., events) by analyzing social media data. In this paper, we use the angle of the positive emotion generated by these topics for the users and the magnitude, both reach and time span, in order to better understand what is happening on social media platforms, mainly Twitter - currently rebranded as X.

Sentiment Analysis is a field in Natural Language Processing that analyzes user opinions and emotions from written language [18], [39], [74], while Event Detection deals with analyzing information diffusion in graph networks [26]. Although there is a large volume of work done on Event Detection using social media data and on Sentiment Analysis of this type of content, in the current literature, there is a shortcoming of the approaches that combine the two domains. There are multiple communities that are involved in mining, gathering, and giving some meaning to the vast amount of content generated daily by the users of those platforms, namely the Network Analysis and Natural Language Processing communities. The two communities are using different types of approaches since they have different purposes:

- For the Network Analysis community, the main purpose is developing methods to deal with the spread and mitigation of harmful content using Event Detection. Event Detection is used to detect the impact and spread of topics on Social Networks using multiple types of approaches such as sliding windows, topic detection, etc.
- For the Natural Language Processing community, the main purpose is to develop methods to analyze the opinion of the users using Sentiment Analysis. Sentiment Analysis uses supervised and unsupervised learning to determine the polarity of textual data. The dimensions of the actual classification of a text are different for each method, ranging from binary (positive/negative or neutral if we cannot tell) to multi-dimensional approaches (happiness, contempt, surprise, and so on).

In this paper, we want to address this shortcoming and bridge the gap between the Network Analysis and Natural Language Processing communities by providing an Event Detection Sentiment Analysis Ensemble Architecture, i.e., EDSA-Ensemble. EDSA-Ensemble first detects events from social media in almost real-time and then applies sentiment analysis on the resulting

topic of each event using an ensemble approach. By combining Network Analysis approaches with the detection of event sentiments using Machine and Deep Learning in Online Social Networks, our architecture brings together the research of these two isolated communities, i.e., Network Analysis and Natural Language Processing. We aim to evaluate individually various types of Event Detection strategies and Sentiment Analysis models and then combine them into an ensemble model to facilitate the detection and spread of events in a content-aware and network-aware manner. By performing such an extensive analysis, we prove the efficiency of our solution in comparison with current state-of-the-art sentiment analysis models (i.e., EDSA-Ensemble outperforms other ensemble models such as [63], [65], [66], [67]) and we manage to fill an existing research gap, i.e., sentiment analysis of Social Media events.

The main research questions we want to address with this work are:

- (Q_1) How can we determine the different events discussed on social media for a given time span?
- (Q_2) Can we accurately determine the individual opinions of users for a given event?
- (Q_3) Can we accurately determine the overall sentiment of these events?

To address these questions, we propose EDSA-Ensemble. Firstly, EDSA-Ensemble uses multiple Event Detection algorithms to determine the important topics that are discussed online. By using multiple algorithms that use different heuristics to detect topics of importance, we ensure that we do not miss any events for a given time span. Secondly, EDSA-Ensemble uses multiple algorithms to determine the users' opinions by analyzing individual tweets. Finally, EDSA-Ensemble determines the overall sentiment of events by aggregating through a voting process the results of the multiple individual sentiment analysis models, thus improving the accuracy of detecting the overall sentiment of an event.

The main objectives of our novel architecture are:

- (O_1) Propose a modular design that determines events on social media using multiple algorithms.
- (O_2) Determine the individual user opinion regarding an event by developing multiple sentiment detection models.
- (O_3) Propose a voting-based ensemble model that determines with high accuracy the overall sentiment of an event.

Thus, the main contributions are as follows:

- (C_1) We propose EDSA-Ensemble: a modular architecture that incorporates multiple Event Detection algorithms and Sentiment Analysis models to accurately detect the important events discussed on social media and their polarity.
- (C_2) We present an in-depth analysis of each module of the proposed architecture on a real-world dataset.
- (C_3) We do a thorough discussion on the obtained experimental results.

This paper is structured as follows. Section II presents a survey of the current state-of-the-art methods for Event Detection and Sentiment Analysis. Section III describes the EDSA-Ensemble

architecture and its main functionalities. Also, we briefly describe how each component interacts with the overall platform and how each chosen algorithm works. Section IV showcases the obtained results, analyzes them, and gives some directions for future trials using the chosen Event Detection and Sentiment Analysis methods. Section VI concludes and presents several new directions and improvements for the proposed solution.

II. RELATED WORK

In this section, we present an overview of the current state of the art related to Event Detection and Sentiment Analysis.

A. Event Detection

Event Detection tasks use a small amount of training data which leads to poor results for unseen/sparsely labeled trigger words as the models are prone to overfitting due to densely labeled trigger words. Thus, Tong et al. [69] propose a novel Enrichment Knowledge Distillation (EKD) model for leveraging external open-domain trigger knowledge to reduce the in-built biases to frequent trigger words in annotations which improves accuracy, especially for unseen or infrequent trigger words.

For social media event detection, current methods learn limited amounts of knowledge and ignore rich semantics and structural information. To address this issue, Cao et al. [11] propose a novel Knowledge-Preserving Incremental Heterogeneous Graph Neural Network (KPGNN) that leverages this information using Graph Neural Networks. KPGNN needs no feature engineering and requires a small amount of parameter tuning. EDGNN [22], a new event detection method using Graph Neural Networks (GNNs) and BERT, leverages topic modeling to enrich the semantics of short texts. To address the problem of collectively detecting multiple events, particularly in cross-sentence settings, Lou et al. [36] encodes semantic information and uses a model for event inter-dependency at a document level that captures both event meaning and connections between them. Furthermore, structural relationships among users, such as online network communities, is rarely used when performing event detection. Thus, Ansah et al. [4] propose SensorTree, a framework that tracks information spread in Twitter communities to identify event spikes. Moreover, most event detection frameworks do not capture the embedded semantic meanings or relationships. To address this, Abebe et al. [1] introduce SEDDaL (Social-based Event Detection, Description, and Linkage framework) which uses heterogeneous sources (e.g., Flickr, YouTube, and Twitter) as input to produce semantically meaningful events interconnected with spatial, temporal, and semantic relationships. Additionally, event detection methods are focused on statistical and syntactical data features, missing deeper textual meaning, i.e., the underlying semantics. Hetiarachchi et al. [28] address this with Embed2Detect, which combines word embedding and hierarchical agglomerative clustering for social media event detection.

Multi-source event propagation involves sharing and collaboration over the same topic from multiple people and there is not a lot of talk on this topic in the current literature. Shi et al. [59]

create a framework that uses previous knowledge about the event in order to obtain a multi-source events propagation model based on individual interest. This solution is not compared with other approaches.

Event detection traditionally relies on using topic modeling and topic labeling techniques [70], [72] on sentence-level information. However, events often unfold across sentences. To address this challenge, Yao et al. [78] propose FGCAN, a Filter-based Gated Contextual Attention Network model, which is augmented with hierarchical contextualized representations to utilize both sentence-level and document-level information.

Missing corpus labels or the limited performance of the classifier can lead to unreliable sentiment time series for event detection. To address this, Wu et al. [76] propose a new method to calibrate the series for event detection based on the evaluation of a sampling dataset.

Information diffusion is the field that analyses how the states of individual nodes in a graph evolve over time [23]. In the case of Online Social Networks, it analyzes how "bursty" topics (i.e., topics that spread information better than the average) gain traction and fade over time, helping us understand user behavior and information flow [17]. Guille et al. [27] survey the state-of-the-art methods for Event Detection and ways of building a taxonomy with good performance (time spread and magnitude). Their analysis focuses on Twitter and detects bursty topics focusing on the evolution in time of the topics. Experimental results prove that the information is diffused by users that appear as the node for which the graph partitions [25]. The analysis includes multiple algorithms and methods, e.g., Peaky Topics [58] that uses normalized term frequency, TSTE (Temporal and Social Terms Evaluation) [12] that uses a five-step approach considering both temporal and social properties, MACD (Moving Average Convergence Divergence) [37] that uses the trend momentum of a topic, SDNML (Sequentially Discounting Normalized Maximum Likelihood) [64] specialized in tweets that also have media and URLs, and OLDA (Online Latent Dirichlet Allocation) [3] which is an improvement to the LDA adding, with the online phase. Guille and Hacid [26] propose an ML-based approach that uses inference of time-dependent diffusion probabilities from a multidimensional analysis of individual behaviors. Guille and Favre [24] propose MABED (Mention-Anomaly-Based Event Detection), a method for event detection in social networks using user mentions MABED builds communities around users and detects events based on mention statistics. It outperforms baselines, i.e., PS (Peakiness Score) [58] and EDCoW (Event Detection with Clustering of Wavelet-based Signals) [75], in topic clarity and temporal accuracy. At the other end of the spectrum, there are time-frame-based Event Detection methods such as Peaky Topics [58], which splits the corpus into equally distributed bins and detects the bursty topics in each bin.

B. Sentiment Analysis

There are two main approaches for Sentiment Analysis: 1) Lexicon Based (unsupervised methods that use word polarity to classify textual data), and 2) Machine Learning (supervised methods that use polarity labeled dataset to build a model).

Sirbu et al. [62] present a comprehensive study of the two main Sentiment Analysis approaches. The authors find that using Multivariate Analysis of Variance (MANOVA) for feature extraction leads to better results. This method effectively captures relationships between sentiment and influencing factors. Thelwall et al. [68] try to determine the opinion presented in Twitter events during peak hours. The results show that events with negative sentiments are more frequent, whereas positive predictions were accurate only during peak hours. Kouloumpis et al. [32] explore combining sentiment lexicons with linguistic features (n-grams, part-of-speech) and metadata (hashtags, emoticons) to improve Twitter sentiment analysis. Pak and Paroubek [45] address the impact of adjective overuse on sentiment analysis by incorporating adjective polarity into a Naive Bayes model, improving sentence-level sentiment detection accuracy.

Deep learning approaches using Convolutional Neural Networks (CNN), Recurrent Neural Networks (RNN), Gated Recurrent Units (GRU), and Long Short-Term Memory (LSTM) are also used for sentiment analysis [31]. Recent studies explore feature ensembles for fuzzy sentiment [53] and Bidirectional emotional Recurrent Units for capturing context in conversations [34], achieving improved sentiment detection on tweets. Yang et al. [77] introduce the SLCABG model, a deep learning model that employs sentiment lexicons and a deep learning architecture with CNN and attention-based Bidirectional GRU (BiGRU) to determine the sentiments from reviews. Onan [43] proposes a CNN-LSTM architecture with a TFIDF weighted GloVe [49] embedding, while Ruz et al. [56] use a Bayesian network classifiers for sentiment analysis. Behera et al. [8] also combines CNN and LSTM to extract the opinion from consumer reviews posted on social media, while Onan [44] combines Bidirectional convolutional recurrent neural with group-wise enhancement mechanism for sentiment analysis. Naseem et al. [42] introduce DICE, a transformer-based sentiment analysis method that improves tweet quality by removing noise and leverages contextual information for sentiment detection. Ensemble models [67] that employ word embeddings [63] and transformers [65], [66] have been successfully used to preserve context when performing sentiment analysis tasks.

Another research direction is aspect-based sentiment analysis (ABSA) which tries to perform a multi-faceted analysis of textual data in order to determine the sentiment regarding both target items as well as target categories [5]. Basiri et al. [6] propose ABCDM, an Attention-based Bidirectional CNN-RNN Deep Model ABCDM. Peng et al. [48] propose an aspect sentiment triplet extraction deep neural model for aspect-based sentiment analysis that extracts *What* (aspect), *How* (sentiment), and *Why* (explanation) from the input text.

C. Event Sentiment Analysis

In the current literature, event sentiment analysis has not yet been widely explored, although it is a very important topic [19].

Fukuhara et al. [21] propose using two graphs to track sentiment and topic trends over time in news articles. One graph shows how sentiment changes for a specific topic, while the other shows how topics with a certain sentiment evolve. Zhou

et al. [81] propose a Tweets Sentiment Analysis Model (TSAM) that extracts the societal interest and users' opinions regarding a social event. TSAM uses a lexicon to extract the individual polarity and, after preprocessing the text and identifying the entities, an overall sentiment score is calculated. Makrehchi et al. [38] propose a new social media text labeling approach that considers stock market events on Twitter. The authors train a classification model using individual tweets to predict the polarity and then aggregate the results to obtain a net sentiment for each day. Patil and Chavan [47] propose SegAnalysis, a framework that performs tweet segmentation, event detection, and sentiment analysis. SegAnalysis uses Naïve Bayes and online clustering to detect events and a lexicon approach to compute the sentiment score. Petrescu et al. [51] propose the use of Logistic Regression and Support Vector Machines to analyze individual events extracted by MABED and OLDA. Basu et al. [7] introduce a new feature set selection framework. The feature set uses Information Gain to extract topics from live-streamed tweets. Bag-of-words with N-Gram identification is applied to extract event features from the textual data as well as part-of-speech linguistic annotation. A decision tree is used to extract the tweets that belong to a topic from the live Twitter stream and their polarity. Zhang et al. [80] propose an End-to-End Event-level Sentiment Analysis (E3SA) approach to solve the task of structured event-level sentiment analysis. To enhance the event-level polarity detection, the method explicitly extracts and models the event structure information, i.e., subject, object, time, and location.

III. METHODOLOGY

In this section, we present our proposed solution for accurate event and sentiment detection in social media. This solution offers two classes of tasks: Event Detection and Sentiment Analysis tasks. For the Event Detection task, we are using a raw frequency-based embedding as a method to vectorize the text. For the Sentiment Analysis task, we are using the standard architectures for the chosen methods.

A. Proposed Architecture

The architecture of the proposed solution EDSA-Ensemble is presented in Fig. 1. Its main functionality and modules are as follows:

- *Twitter Collection Module* - is constantly listening for Twitter data in order to fetch and pass it on to the Preprocessing Module.
- *Preprocessing Module* - applies different models for text representation on the received datasets from the Collection Module.
- *Event Detection Module* - is used to detect bursty topics either on cleaned or raw text.
- *Sentiment Analysis Module* - applies different methods for sentiment analysis on the preprocessed datasets, e.g., Classical ML Models (Naïve Bayes, Support Vector Machines) and Deep Models (BiLSTM and Transformers), and uses an ensemble voting system to determine the correct class for each document.

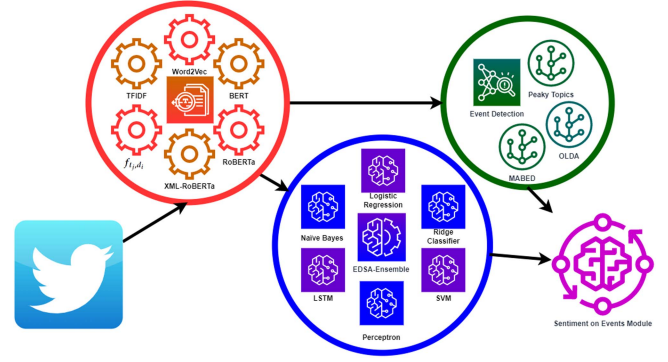


Fig. 1. EDSA-Ensemble architecture diagram.

- *Sentiments on Events Module* - applies Sentiment Analysis methods on the outputs of the Event Detection Module in order to determine if an event has a positive or negative impact over the network.

In the following subsections, we present a detailed description of each architectural module from Fig. 1.

B. Preprocessing Module

In Natural Language Processing, words are represented using numerical values or weight. Thus, a weight (numerical value or vector) is assigned to a word. These weights can encapsulate the importance of a word in a corpus of text and also its context. We employ three models for representing terms from the corpus (1) Bag-of-Words model using raw frequency f_{t_j, d_i} and $TFIDF$; (2) Word Embedding model using Word2Vec; and (3) Transformer Embedding model using BERT, RoBERTa, and XLM-RoBERTa.

1) *Bag-of-Words Model*: The Bag-of-Words model can be built using multiple weighting schemes. Given a corpus of documents $D = \{d_i \mid i \in \overline{1, n}\}$ ($n = ||D||$ is the size of the documents) and a vocabulary $V = \{t_j \mid j \in \overline{1, m}\}$ ($m = ||V||$ is the size of the vocabulary) that contains the set of unique terms t_j that appear in dataset, we can define using the Bag-of-Words model multiple weighting schemes that encode documents and terms into vectors. EDSA employ two of them: i) the raw frequency and (TF) ; ii) term frequency-inverse document frequency $(TFIDF)$.

The raw frequency (f_{t_j, d_i}) counts the number of co-occurrences of a word in a document. It provides no context and adds bias toward longer documents. The term frequency-inverse document frequency $(TFIDF)$ weighting scheme is often used as a central tool for scoring and ranking terms in a document. The $TFIDF$ weight (1) is computed by multiplying the augmented term frequency (TF) with the inverse document frequency (IDF) . For $TFIDF$, the importance increases proportionally to the number of times a term appears in the document but is offset by the frequency of the term in the corpus. To prevent a bias towards long documents, TF (2) is computed as the number of co-occurrences f_{t_j, d_i} of a word in a document normalized by the length of the document ($\sum_{t' \in d} f_{t', d}$) when computing the TF [46]. IDF (3) is used for minimizing the importance of

common terms that bring no information value by reducing the TF weight by a factor that grows with the collection frequency n_j of a term t_j .

$$TFIDF(t_j, d_i, D) = TF(t_j, d_i) \cdot IDF(t_j, D) \quad (1)$$

$$TF(t_j, d_i) = \frac{f_{t_j, d_i}}{\sum_{t' \in d} f_{t', d_i}} \quad (2)$$

$$IDF(t_j, D) = \log \frac{n}{n_j} \quad (3)$$

Using the term weights, we can construct a document-term matrix $A = \{w_{ij} \mid i = \overline{1, n} \wedge j = \overline{1, m}\}$, where rows correspond to documents and terms to columns. The cell value w_{ij} is the weight (e.g., f_{t_j, d_i} , TF, TFIDF, etc.) of term t_j in document d_i .

2) *Word Embedding Model*: Word embedding consists of a set of techniques used for weighting word and text analysis. It features learning techniques that assign numerical values (usually as a vector) to linguistic words. The proposed solution uses the Word2Vec technique to construct a Common Bag Of Words (CBOW) model [40], which is a type of word embedding.

The CBOW model tries to predict the target word by looking at the words situated at its left and right sides. The model uses one-hot encoder vectors of size m as input context $(x_{-m}, \dots, x_{-1}, x_1, \dots, x_m)$ for each target word $x_i \in R^m$. The embedding word vectors for the context are $v_i = Vx_i$, where V is the input word matrix and $i \in \{-m, \dots, -1, 1, \dots, m\}$.

The dot product of vectors is used to push similar words close to each other and maximize the model's performance. The output score is computed using the average of all word embedding vectors as $Z = U\hat{v}$, where U is the output word matrix and $\hat{v}$ is the average of m context vectors for the target word target [61]. The output is a word vector ($\hat{y}$) obtained from the softmax function (4).

$$\hat{y} = \frac{e^{\hat{y}}}{\sum_{k=1}^{|V|} e^{\hat{y}_k}} \quad (4)$$

3) *Transformers*: In machine learning, a transformer is a deep learning model that adopts the mechanism of self-attention which balances the importance of each token in the input data. Transformers were introduced in 2017 by a team at Google Brain and they became a popular approach in NLP tasks, replacing RNN models such as long short-term memory (LSTM), obtaining state-of-the-art performance with the rise of BERT [16]. Other popular approaches such as RoBERTa [35] and XLM-RoBERTa [14] have the same architecture but have different training data which result in a different performance for each task that they are not specialized in.

We are going to benchmark the previously mentioned transformers against our Twitter data since transformers obtained state-of-the-art performance in NLP tasks and we are using multiple variations in order to see how the increasing quantity of data, RoBERTa, affects the base, BERT, or how the multilingual features, XLM-RoBERTa, do.

a) *BERT*: BERT (Bidirectional Encoder Representations from Transformers) [16] proves the importance of bidirectional

pre-training by surpassing the previous state-of-the-art approach which was a left-to-right approach, while BERT instead uses the self-attention mechanism to unify these two stages, as encoding a concatenated text pair with self-attention effectively includes bidirectional cross attention between two sentences. For this paper, the pre-trained embeddings used are the Word-Piece embeddings that contain a vocabulary of 30,000 tokens and a pre-training on the following datasets (1) BooksCorpus (800 M words) and (2) English Wikipedia (2,500 M words) using everything except lists, tables, and headers.

The core BERT framework has 2 major steps:

- 1) Pre-training BERT: Masked LM bidirectional training with randomly masked tokens (15% of the input) trying to predict those tokens, the only downside being that the mask is not present during the fine-tuning. In order to mitigate this for the randomly chosen tokens 80% of the time the swap is done with the mask, 10% with another random token and 10% of the time it is left as it is. Next Sentence Prediction is used for Q&A modelling where 50% of the time the sequence is correct.
- 2) Fine-tuning: A straightforward strategy since the self-attention mechanism in the Transformer allows BERT to model many downstream tasks.

Devlin et al. [16] state that it is critical to use a document-level corpus rather than a shuffled sentence-level corpus such as the Billion Word Benchmark [13] in order to extract long contiguous sequences.

b) *RoBERTa*: RoBERTa (Robustly Optimized BERT Pre-training Approach) [35] is a retraining of BERT with improved training methodology, more data, and more computing power. To improve the training procedure, RoBERTa removes the Next Sentence Prediction (NSP) task from BERT's pre-training and introduces dynamic masking so that the masked token changes during the training epochs. Larger batch-training sizes were also found to be more useful in the training procedure. Importantly, RoBERTa uses 160 GB of text for pre-training, including 16 GB of Books Corpus and English Wikipedia used in BERT. The additional data included CommonCrawl News dataset (63 million articles, 76 GB), Web text corpus (38 GB), and Stories from Common Crawl (31 GB). This coupled with a whopping 1024 V100 Tesla GPUs running for a day, led to pre-training of RoBERTa.

c) *XLM-RoBERTa*: XLM-RoBERTa [14] model is a transformer-based pre-trained multilingual masked language model able to determine the correct language just from the input ids, being pre-trained on 2.5 TB of filtered CommonCrawl data containing 100 languages, mainly being a multilingual version of RoBERTa.

Due to the fact that we are using social data for this task, we might encounter multilingual content even if the selected language was English and for this particular reason, another transformer seems like a good benchmark.

C. Event Detection Module

For Event Detection, we use the Mention-Anomaly-Based Event Detection (MABED) [24], Online LDA (OLDA) [3],

and Peak Topics [58] algorithms. All approaches detect bursty topics either on cleaned or raw text. The cleaned text results by removing the stop words and any non-character from the raw text and applying lemmatization.

1) *Mention-Anomaly-Based Event Detection: Mention-Anomaly-Based Event Detection (MABED)* is an efficient statistical method used for detecting events in social networks since it is immune to social media bias consisting of unrelated texts on a given topic [24]. Although social networks can contain spam messages with no actual intent posted around certain hours, this method is efficient when filtering irrelevant content. MABED is not always viable since, without external context, it can sometimes distort some types of events.

To identify a bursty topic and detect an event for a period of time $I = [a; b]$ and a main word t , a weight $w_{t,q}$ is computed for each candidate word t'_q in the time slice i (5). The weight $w_{t,q}$ is in the range $[0, 1]$. The higher the weight, the more important the candidate term t'_q is to the event. The affine function $\rho_{O_{t,t'_q}}$ (6) is used to compute the weight. This function corresponds to the first order auto-correlation of the time series for N_t^i (number of tweets in the time-slice i that contain the main word t) and $N_{t'_q}^i$ (number of tweets in the time-slice i that contain the candidate word t'_q) [20].

$$w_{t,q} = \frac{\rho_{O_{t,t'_q}} + 1}{2} \quad (5)$$

$$\rho_{O_{t,t'_q}} = \frac{\sum_{i=a+1}^b A_{t,t'_q}}{(b-a-1)A_t A_{t'_q}} \quad (6)$$

Where:

$$\begin{aligned} 1) A_{t,t'_q} &= (N_t^i - N_t^{i-1})(N_{t'_q}^i - N_{t'_q}^{i-1}) \\ 2) A_t^2 &= \frac{\sum_{i=a+1}^b (N_t^i - N_t^{i-1})^2}{(b-a-1)}; \\ 3) A_{t'_q}^2 &= \frac{\sum_{i=a+1}^b (N_{t'_q}^i - N_{t'_q}^{i-1})^2}{(b-a-1)}; \end{aligned}$$

2) *Online LDA*: Online LDA [3], is an update on Latent Dirichlet Allocation [9] which uses a non-Markov on-line LDA Gibbs sampler topic model. This model is capable of detecting bursty topics by capturing thematic patterns and identifying emerging topics and their changes over time. This approach is sometimes better than the original LDA approach, as it updates the previous model at each time frame.

OLDA is a hierarchical Bayesian network that relates words and documents through latent topics using a probabilistic formula for a given term to be assigned to a given topic in the (current) context (7).

$$\begin{aligned} P(z_i = j \mid z_{-i}, w_{d_i}, \alpha, \beta) &\propto \frac{C_{w_{-i},j}^{VK} + \beta_{w_{i,j}}}{\sum_{v=1}^V C_{v_{-i},j}^{VK} + \beta_{v,j}} \\ &\times \frac{C_{d_{-i},j}^{DK} + \alpha_{w_{i,j}}}{\sum_{v=1}^V C_{d_{-i},j}^{DK} + \alpha_{d,j}} \quad (7) \end{aligned}$$

Where:

- D is the total number of documents;
- K is the number of topics;
- V is the total number of unique words;

- $C_{w_{-i},j}^{VK}$ is the number of times word w is assigned to topic j , not including the current token instance i
- $C_{d_{-i},j}^{DK}$ is the number of times topic j is assigned to some word token in document d , not including the current instance i ;
- z_{-i} are all other word tokens;
- w_{d_i} are the unique word associated with the i -th token in document d at the current time;
- β_k is the V vector of priors for topic k at the current time;
- α_k is the K vector of priors for document k at the current time.

3) *Peak Topics*: Identifying peaks in a time series is a very important task in many cases when immediate action is required which may point out 1) high demands in power distribution data, 2) increased CPU utilization, 3) burst in traffic or 4) hasty price increasing in financial data.

For the particular case of stream-series on Twitter, let us consider $S = \{t_1, t_2, \dots, t_M\}$ a text stream collected during a time period. To apply the Peak Topics algorithm, this stream should be divided in multiple equally distributed time-frame based bins [58]. After the splitting, we obtain $T = \{x_1, x_2, \dots, x_N\}$ where $x_i \in T$ represents a tweet posted in the time-frame T . The size of the bin, time-window, should be adapted relative to the total duration of the initial text stream.

For every x_i value in the T , a check value is calculated and is used to determine if a point represents a local maximum or not. To determine the spikes in a given frame, we consider many small sub-bins. For each sub-bin we run a spike detection algorithm based on the number of tweets in those sub-bins.

Peak Topics performs better at detecting bursty topics in narrow time frames compared to MABED and OLDA, but it does not take into consideration the evolution in time of the topic [17].

D. Sentiment Analysis Module

For Sentiment Analysis, we use Naïve Bayes (NB), Logistic Regression (LR), Ridge Classifier (RC), Support Vector Machines (SVM), and Long Short-Term Memory Neural Networks (LSTM). All approaches behave well on both raw texts or with feature engineering, and negation fusion. Each method has an advantage over the other either in run-time, particular benchmarks, or overall benchmarks.

1) *Naïve Bayes*: Naïve Bayes (NB) is one of the base methods for sentiment analysis with good results most of the time, in our paper is the baseline. NB is a probabilistic classification algorithm that computes the probability $p(y = y_\kappa \mid d = d_i) = p(y = y_\kappa \mid t_1 = t_{i1}, \dots, t_m = t_{im})$ of a document $d_i = \{t_{ij} \mid j = \overline{1, m}\}$ to a given class $y_\kappa \in C$ ($i = \overline{1, n}$) using the Bayes Theorem. Equation (8) presents the Bayes Theorem, where $p(d = d_i)$ and $p(y = y_\kappa)$ are the probability of a document, respectively a class, and $p(d = d_i \mid y = y_\kappa) = p(t_1 = t_{i1}, \dots, t_m = t_{im} \mid y = y_\kappa)$ is the probability of class y_κ given a document d_i .

$$p(y = y_\kappa \mid t_1 = t_{i1}, \dots, t_m = t_{im})$$

$$= \frac{p(y = y_\kappa) p(t_1 = t_{i1}, \dots, t_m = t_{im} \mid y = y_\kappa)}{p(d = d_i)} \quad (8)$$

The denominator $p(d = d_i)$ is constant and the numerator is equivalent to the joint probability $p(y = y_\kappa, t_1 = t_{i1}, \dots, t_m = t_{im})$.

Furthermore, all the terms are conditionally independent given a class y_κ , thus $p(y = y_\kappa, t_1 = t_{i1}, \dots, t_m = t_{im}) = p(y = y_\kappa) \prod_{j=1}^m p(t_j = t_{ij} \mid y = y_\kappa)$ and the model $p(y = y_\kappa \mid t_1 = t_{i1}, \dots, t_m = t_{im}) \propto p(y = y_\kappa, t_1 = t_{i1}, \dots, t_m = t_{im})$.

The Naïve Bayes classifier tries to estimate the class $\hat{y}_i$ for a document d_i using (9).

$$\hat{y}_i = \operatorname{argmax}_{y_\kappa \in C} p(y = y_\kappa) \prod_{j=1}^m p(t_j = t_{ij} \mid y = y_\kappa) \quad (9)$$

Multinomial Naïve Bayes uses a multinomial representation to model the distribution of words in a document. Thus, this model uses the distribution of probabilities to determine the class where each word appears frequently (10). The model makes two assumptions [55]: 1) a document is a sequence of words and 2) the position of a word is generated independently. The model uses the co-occurrence f_{t_j, d_i} of a term t_j in the document d_i . Equation (11) presents the Multinomial Naïve Bayes model, while (12) presents the estimated values $\hat{y}_i$.

$$p(t_1 = t_{i1}, \dots, t_m = t_{im} \mid y = y_\kappa) = \frac{(\sum_{j=1}^m f_{t_j, d_i})!}{\prod_{j=1}^m f_{t_j, d_i}!} \prod_{j=1}^m p(t_j = t_{ij} \mid y = y_\kappa)^{f_{t_j, d_i}} \quad (10)$$

$$p(y = y_\kappa \mid t_1 = t_{i1}, \dots, t_m = t_{im}) \propto p(y = y_\kappa) \frac{(\sum_{j=1}^m f_{t_j, d_i})!}{\prod_{j=1}^m f_{t_j, d_i}!} \prod_{j=1}^m p(t_j = t_{ij} \mid y = y_\kappa)^{f_{t_j, d_i}} \quad (11)$$

$$\hat{y}_i = \operatorname{argmax}_{y_\kappa \in C} p(y_i = y_\kappa) \frac{(\sum_{j=1}^m f_{t_j, d_i})!}{\prod_{j=1}^m f_{t_j, d_i}!} \prod_{j=1}^m p(t_j = t_{ij} \mid y = y_\kappa)^{f_{t_j, d_i}} \quad (12)$$

2) *Logistic Regression*: Logistic regression (LR) [29] is a statistical model, widely used, that in its basic form uses a logistic function to model a binary classification for the given dataset. It tries to build a function, a linear one for this case that summarizes the input data. Logistic regression is a supervised method that behaves better than Linear Regression [41] when applied to a sparse dataset, and builds classification models within a probabilistic context [2]. We consider that a dataset is sparse when the matrix of its representation has a lot of zeros and a few one values. If we represent the words depending on their appearance into a certain block of text (tweet) or not, we obtain a sparse matrix since the tweets have at most 140 characters.

Given a set of classes $C = \{y_\kappa \mid \kappa = \overline{1, k}\}$ and a document dataset $D = \{d_i \mid i = \overline{1, n}\}$ described by the document-term matrix A , the model computes the probability of a class y_κ given a document d_i as $p(y = y_\kappa \mid d = d_i)$. The probability is

the sigmoid function that maps documents to classes in order to determine the parameters of the vector $\beta = \{\beta_0, \beta_1, \dots, \beta_m\}$ that fit the regression line. Equation (13) presents the Logistic Regression probability, where $x_i = [1] \oplus a_i$ is a vector composed by concatenating (operator $\oplus$) a one element vector with the value 1 to the line a_i to accommodate the β_0 . The vector a_i is the corresponding line in matrix A to document d_i . The one-element vector is used to accommodate the intercept of the regression line, i.e., β_0 .

$$p(y = y_\kappa \mid d = d_i) = \frac{1}{1 + e^{-\beta^T x_i}} \quad (13)$$

To estimate the classes $\hat{y}_i = \operatorname{argmax}_{y_\kappa \in C} p(y = y_\kappa \mid d = d_i)$ and to determine the parameters β that give the best results, the algorithm maximises the log-likelihood using gradient descent [79]. The log-likelihood function guarantees that the gradient descent algorithm can converge to the global minimum. Equation (14) presents the log-likelihood function $l(\beta)$, where y_i is the real class of document d_i .

$$l(\beta) = \sum_{i=1}^n (y_i x_i \beta - \log(1 + e^{x_i \beta})) \quad (14)$$

3) *Ridge Classifier*: The Ridge Classifier (RC) [50], [60] or the Ridge regression is a learning approach that first normalizes the target values in the range $\{-1, 1\}$ and then treats the problem as a regression task (multi-output regression in the multi-class case). Experimentally this method gives results at least as good as the Logistic Regression. The cost function of the Ridge Regression is adjusted by adding a penalty λ comparable to the square of the magnitude of the coefficients. Equation (15) presents the log-likelihood for the Ridge Regression.

$$l(\beta) = \sum_{i=1}^n (y_i x_i \beta - \log(1 + e^{x_i \beta})) - \lambda \sum_{j=0}^m \beta_j^2 \quad (15)$$

4) *Support Vector Machines*: Support vector machines (SVM) [15] is a supervised learning algorithm used in ML for classifications. Given a tagged dataset, SVM builds a function that separates the space of this given dataset into two classes. In Sentiment Analysis, this type of model is used for binary classification, where the classes are positive and negative representing the tweets' sentiments. For the Sentiment Analysis tasks, we employ a linear SVM. The SVM estimates the class $\hat{y}_i = \operatorname{argmax}_{y_\kappa \in C} f(x_i, y_\kappa)$. Equation (16) is used to separate the values of the training set using into the positive (+1) and negative classes (-1), where w_i^T is the weight vector and b is the bias.

$$f(x_i, y) = \operatorname{sign}(w_i^T x_i + b) \quad (16)$$

5) *Perceptron*: A perceptron is a type of linear classifier used in neural networks usually in a final layer such as softmax which uses the normalized exponential function, (18), as an activation function to distinguish between the classes of the binary classification. The standard representation is given by (17), where $x_t \in \mathbb{R}^m$ is the input of the perceptron with each dimension being a feature of the given input and $y_t = f(x_t)$ being the class, where w is a vector of real-valued weights,

$w \cdot x_t$ is the dot product $\sum_{i=1}^m w_i x_{ti}$, b is the bias, and $\sigma(z)$ is the activation function.

$$\hat{y}_t = f(x_t) = \sigma(wx_t + b) \quad (17)$$

The input of the perceptron is a vector that is weighted in the activation function in order to produce an output in the $[0, 1]$ domain usually being combined with a threshold to give the predicted class $\hat{y}_t$ for the input x_t . In our architectures, we are using the softmax activation function (18) as the final layer for predicting the class $\hat{y}_t$.

$$\sigma(x_t) = \frac{e^{w x_t}}{\sum_{j=1}^m e^{w x_{tj}}} \quad (18)$$

6) *Long Short-Term Memory*: A Neural Network based approach for text classification with good results is Long short-term memory (LSTM) [57]. LSTM is an artificial recurrent neural network (RNN) architecture used in the field of deep learning. Unlike standard feed-forward neural networks, an LSTM has feedback connections. It can not only process single data points so it has been adapted to work for sentiment analysis.

The LSTM uses two state components for learning the current input vector: (1) a short-term memory component represented by the hidden state, and (2) a long-term memory component represented by the internal cell state. The cell also incorporates a gating mechanism with three gates: (1) input gate ($i_t \in \mathbb{R}^n$), (2) forget gate ($f_t \in \mathbb{R}^n$), and (3) output gate ($o_t \in \mathbb{R}^n$). The cell state corresponds to the long-term memory. LSTM avoids the vanishing and the exploding gradient issues. Equation (19) presents the compact forms for the state updates of the LSTM unit for a time step t , where:

- $x_t \in \mathbb{R}^m$ is the input features vector of dimension m ;
- $h_t \in \mathbb{R}^n$ is the hidden state vector as well as the unit's output vector of dimension n , where the initial value is $h_0 = 0$;
- $\tilde{c}_t \in \mathbb{R}^n$ is the input activation vector;
- $c_t \in \mathbb{R}^n$ is the cell state vector, with the initial value $c_0 = 0$;
- $W_i, W_f, W_o, W_c \in \mathbb{R}^{n \times m}$ are the weight matrices corresponding to the current input of the input gate, output gate, forget gate, and the cell state;
- $V_i, V_f, V_o, V_c \in \mathbb{R}^{n \times n}$ are the weight matrices corresponding to the hidden output of the previous state for the current input of the input gate, output gate, forget gate, and the cell state;
- $b_i, b_f, b_o, b_c \in \mathbb{R}^n$ are the bias vectors corresponding to the current input of the input gate, output gate, forget gate, and the cell state;
- $\delta_s(z) = \sigma_g(z) = \frac{1}{1+e^{-z}} \in [0, 1]$ is the sigmoid activation function;
- $\delta_h(z) = \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}} \in [-1, 1]$ is the hyperbolic tangent activation function;
- $\odot$ is the element-wise product, i.e., Hadamard Product.

$$i_t = \delta_s(W_i x_t + V_i h_{t-1} + b_i)$$

$$f_t = \delta_s(W_f x_t + V_f h_{t-1} + b_f)$$

$$o_t = \delta_s(W_o x_t + V_o h_{t-1} + b_o)$$

$$\tilde{c}_t = \delta_h(W_c x_t + V_c h_{t-1} + b_c)$$

$$c_t = i_t \odot \tilde{c}_t + f_t \odot c_{t-1}$$

$$h_t = o_t \odot \delta_h(c_t) \quad (19)$$

E. EDSA-Ensemble: Event Detection Sentiment Analysis Ensemble

The Event Detection methods produce a list of event keywords. Each of these lists encapsulates the overall topic. To each topic, a list with all the tweets discussing it is also attached. By applying Sentiment Analysis models over the list of tweets belonging to a topic, then we can determine the positive or negative impact of the event over the network.

By combining both the event magnitude, i.e., the importance of the event, its lifespan, and the event's polarity, we can better understand user behavior and how the event's topic impacts and the attention given to it influences the overall user response. Thus, we can analyze if topics that have a longer lifespan and higher magnitude influence the users in a positive or negative way. Furthermore, by analyzing individual tweets we can determine opinion-based communities.

The proposed model EDSA-Ensemble (Event Detection Sentiment Analysis Ensemble Model) is implemented through Algorithm 1. As input, the EDSA-Ensemble receives a set of documents D . The output of the model is (1) an individual-sentiment set I_S that contains the individual sentiment of each user for each event obtained by each Event Detection method *ED-Methods* and (2) an event-sentiment set E_S that contains the sentiment obtained for each event obtained by each Event Detection method *ED-Methods*. The individual and event sentiments are detected using an ensemble of Sentiment Analysis models. Thus, for each event $e_i \in E$ determined by an Event Detection method $ED_{method} \in ED-Methods$, we first compute the polarity obtained by each $SA_{models} \in SA-Models$. Afterward, we compute the overall polarity of an event e_i given by a ED_{method} by taking the most frequent sentiment, i.e., the *Mode*, obtained by the different *SA-Models*. To improve the runtime performance of the algorithm, all the iterations over the *ED-Methods* and *SA-Models* are done in parallel.

The algorithm works as follows. Firstly, we initialize the individual-sentiment set I_S (Line 1) and event-sentiment set E_S (Line 2). This set contains the most frequent sentiment obtained by each *SA-Models* of each event e_j obtained by each *ED-Methods*. Secondly, we iterate over each *ED-Methods* (Line 3). For each ED_{method} , we initialize an event-sentiment set E_S^{method} (Line 5), preprocess the text using the text preprocessing function *EDTP* depended on the for each ED_{method} and extract the preprocess corpus D_C (Line 6), and extract the events for the ED_{method} and store it in the set E (Line 7).

For each event $e_i \in E$ (Line 8), we initialize (1) and individual-sentiment set I_{EDSA} (Line 9) where we store the most frequent sentiment obtained by each of the *SA-Models* and (2) an event-sentiment set E_{EDSA} (Line 10) where we store the most frequent sentiment obtained by each of the *SA-Models* for an event (Line 11). To obtain the sentiment using a model SA_{model} of an event e_i , we initialize an empty list $S^{(i)} = \emptyset$ for storing the sentiment for each document belonging to the

Algorithm 1: EDSA-Ensemble.

```

Input : document set  $D$ 
Output: individual-sentiment list  $I_S$  and event-sentiment set  $E_S$ 
/* initialize the individual-sentiment and event-sentiment sets */
1  $I_S = \emptyset$ ;
2  $E_S = \emptyset$ ;
/* for each event detection algorithm */
3 foreach  $ED_{method} \in ED\text{-}Methods$  do
    /* initialize the individual-sentiment and event-sentiment sets for the current
       event detection method */
4    $I_S^{method} = \emptyset$ ;
5    $E_S^{method} = \emptyset$ ;
    /* event detection text preprocessing */
6    $D_C = EDTP(D, ED_{method})$ ;
    /* extract the events */
7    $E = ED(D_C, ED_{method})$ ;
8   foreach  $e_i \in E$  do
    /* initialize the individual-sentiment and event-sentiment sets for the
       current sentiment analysis model */
9      $I_{EDSA} = \emptyset$ ;
10     $E_{EDSA} = \emptyset$ ;
11    foreach  $SA_{model} \in SA\text{-}Models$  do
    /* initialize the sentiment set  $S^{(i)}$  for event  $e_i$  */
12     $S^{(i)} = \emptyset$ ;
    /* extract the set of documents for event  $e_i$  */
13     $D^{(i)} = extract(e_i, D)$ ;
    /* sentiment analysis text preprocessing */
14     $D_C^{(i)} = SATP(D^{(i)}, SA_{model})$ ;
15    foreach  $d \in D_C^{(i)}$  do
    /* extract polarity for each documents  $d$  */
16     $s = SA(d, SA_{model})$ ;
    /* update the sentiment list */
17     $S^{(i)} = S^{(i)} \cup \{s\}$ ;
    /* extract the most frequent sentiment  $s_f^{(i)}$  for event  $e_i$  */
18     $s_f^{(i)} = Mode(S^{(i)}, SA_{model})$ ;
    /* update  $E_{EDSA}$  with the sentiment  $s_f^{(i)}$  of event  $e_i$  */
19     $I_{EDSA} = I_{EDSA} \cup \{S^{(i)}\}$ ;
    /* update  $E_{EDSA}$  with the sentiment  $s_f^{(i)}$  of event  $e_i$  */
20     $E_{EDSA} = E_{EDSA} \cup \{(e_i, s_f^{(i)})\}$ ;
    /* update  $I_S^{method}$  with the most frequent sentiment */
21     $I_S^{method} = Mode(I_{EDSA})$ ;
    /* update  $E_S^{method}$  with the most frequent sentiment of event  $e_i$  */
22     $E_S^{method} = Mode(E_{EDSA})$ ;
    /* update  $I_S$  with the most frequent individual-sentiment */
23     $I_S = I_S \cup \{I_S^{method}\}$ ;
    /* update  $E_S$  with the most frequent sentiment of event  $e_i$  */
24     $E_S = E_S \cup \{E_S^{method}\}$ ;
/* return the  $I_S$  and  $E_S$  lists */
25 return  $I_S, E_S$ ;

```

current analyzed event e_i (Line 12). Then, we extract the set of documents $D^{(i)} = extract(e_i, D)$ that belong to the event e_i (Line 13) and then obtain the clean document set $D_C^{(i)} = SATP(D^{(i)})$ by applying the text preprocessing steps for the current SA_{model} (Line 14).

For each document d in set $D_C^{(i)}$, we detect the sentiment $s = SA(d)$ and update the sentiment list $S^{(i)} = S^{(i)} \cup \{s\}$ for the event e_i (Lines 15 and 17). During this step, we extract individual user opinions regarding the event. After these steps, we label the event e_i with its most frequent sentiment

$s_f^{(i)} = \text{Mode}(S^{(i)})$ (Lines 18 and update (1) the individual-sentiment set $I_{EDSA} = I_{EDSA} \cup \{S^i\}$ (Line 19), and (2) the event-sentiment set $E_{EDSA} = E_{EDSA} \cup \{(e_i, s_f^{(i)})\}$ for the current SA_{model} (Line 20)

After we obtain a polarity of an event e_i with each SA_{model} , we update (1) the I_S^{method} with the most frequent individual sentiment $I_S^{method} = \text{Mode}(I_{EDSA})$ (Line 21) and (2) the E_S^{method} with the most frequent sentiment $E_S^{method} = \text{Mode}(E_{EDSA})$ for the event e_i (Line 22). After we obtain the ensemble results, i.e., E_S^{method} and I_S^{method} , we update (1) the individual-sentiment set I_S (Line 23), i.e., a set of pairs $\langle \text{tweet}, \text{sentiment} \rangle$, and (2) the event-sentiment set E_S (Line 24) using a majority voting method, i.e., select the most frequent sentiment for the tweets belonging to the events to obtain the pairs $\langle \text{event}, \text{sentiment} \rangle$.

When we finished the iterations over the *ED-Methods*, the algorithm returns the final individual-sentiment set I_S and event-sentiment set E_S (Line 25).

IV. EXPERIMENTAL RESULTS

In this section, firstly, we present the dataset used for the three tasks: (1) Event Detection, (2) Individual Sentiment Analysis, and (3) Event Sentiment Analysis. Secondly, we discuss the results obtained with the event detection algorithms individually. Thirdly, we analyze the results obtained for the individual user sentiment of events when using each classification model separately and when using EDSA-Ensemble. Finally, we present the result of detecting the overall sentiments of events for each individual model and the EDSA-Ensemble.

A. Dataset

We use Sentiment140 for our experiments. The dataset contains 1.6 million annotated tweets with the labels negative, neutral, and positive. Besides the textual content, the tweets are also timestamped with the date. The date and the textual content are used to extract events, while the textual content and the labels are employed to build the classification models used by EDSA-Ensemble.

To detect events and to determine the impact of text preprocessing on this task, we create three different text preprocessing pipelines: Minimal Text (MT), Pure Text (PT), and Clean Text (CT). After preprocessing, we weight the terms to prepare the text for the event detection algorithms. We employ the raw frequency (f_{t_j, d_i}), i.e., the number of times a term appears in a document, to weight the words before extracting events with MABED and OLDA. For Peak Topics, we used the *TFIDF* weighting scheme. We use the entire corpus for the Event Detection experiments. Table I presents the steps employed to clean the text used by the three different preprocessing pipelines.

For sentiment analysis, we apply two different preprocessing pipelines: Sentiment Clean Text (SCT) and Sentiment Clean Text with Feature (SFE). When using the SCT pipeline, the text is split into tokens and the punctuation is removed. To preserve the sentiment polarity, we keep the case of words, duplicate letters in words, and stop words. When using the SFE pipeline,

TABLE I
EVENT DETECTION TEXT PREPROCESSING

Text Preprocessing	MT	PT	CT
Lowercase the text	✓	✓	✓
Split the text into tokens	✓	✓	✓
Remove punctuation	✓	✓	✓
Remove stop words	x	✓	✓
Extract the lemma	x	x	✓

TABLE II
SENTIMENT ANALYSIS TEXT PREPROCESSING

Text Preprocessing	SCT	SFE
Lowercase the text	x	x
Split the text into tokens	✓	✓
Remove punctuation	✓	✓
Remove stop words	x	x
Extract the lemma	x	✓
Expand contractions	x	✓
Text enhancement	x	✓

TABLE III
DATASETS FOR SENTIMENT ANALYSIS

Dataset	Tweets	SCT features	SFE Features
C_1	20 000	30 651	31 018
C_2	500 000	307 890	311 240
C_3	1 600 233	684 492	691 646

besides splitting the text into tokens and removing punctuation, we apply feature engineering. First, we expand the contractions, e.g., “don’t” becomes “do not”. Second, we enhance the text by combining the negation with the words that follow, i.e., “not good” becomes “not_good”. Table II presents the steps used to create SCT and SFE.

For Sentiment Analysis, we split the corpus into three subsets to determine the scalability of the Sentiment Analysis models. We create two subsets, i.e., C_1 and C_2 , from the entire dataset, i.e., C_3 , by selecting tweets at random while keeping the label distribution not to over-fit the models. Table III presents the three datasets and the number of features for both preprocessing pipelines.

After applying text preprocessing on the three subsets, i.e., C_1 , C_2 , and C_3 , the words are weighted as follows. We use the raw frequency (f_{t_j, d_i}) for the Naïve Bayes (NB) and Ridge Classifier (RC) models. For the Logistic Regression (LR) and Support Vector Machines (SVM), we encode the text using *TFIDF*. We create word embeddings using Word2Vec to encode the textual data and train the LSTM model. We use transformer embeddings, i.e., BERT, RoBERTa, and XLM-RoBERTa, to train the Perceptron model.

B. Event Detection

For this first set of experiments, we extract 50 events for the time window between April and July 2009. Each extracted event is described by the top-10 most relevant keywords. To determine which preprocessing strategy is better suited to extract clear and human-readable events, we did experiments with all three text preprocessing pipelines, i.e., MT, PT, and CT.

For the first set of experiments, we observe that MABED manages to extract human-readable events when employing CT

TABLE IV
EXAMPLES OF EVENTS OBTAINED WITH MT AND PT USING EACH ED METHOD

Method	Preprocessing	Magnitude	Start Date	End Date	Event Keywords
Peaky Topics	MT	14 342	2009-04-06 22:23:41	2009-06-07 10:38:26	ha http would make love
Peaky Topics	MT	15 733	2009-04-06 22:23:41	2009-06-08 10:29:26	http would ha love lol
Peaky Topics	PT	17 523	2009-04-06 22:27:48	2009-06-09 11:25:21	work ha lol like last
Peaky Topics	PT	17 867	2009-04-06 22:22:48	2009-06-10 12:27:47	im would great lol like
OLDA	MT	12 767	2009-04-06 22:23:11	2009-06-08 13:32:42	quot twitter http new hey
OLDA	MT	16 607	2009-04-06 22:23:23	2009-06-09 12:42:34	sad school back go need
OLDA	PT	20 358	2009-04-06 22:20:20	2009-06-10 15:17:19	na love miss go wish
OLDA	PT	24 245	2009-04-06 22:22:48	2009-06-11 12:19:17	lol http like amp love
MABED	MT	11 256	2009-04-06 22:27:18	2009-06-10 14:18:19	ya nt morning girl quot
MABED	MT	14 214	2009-04-06 22:20:19	2009-06-11 12:24:21	nt wish amp thought watch
MABED	PT	16 497	2009-04-06 22:21:49	2009-06-12 17:19:34	sad nt hate sleep show
MABED	PT	19 466	2009-04-06 22:21:39	2009-06-13 20:34:52	gon nt wan bed sad miss

TABLE V
TOP-10 EVENTS BY MAGNITUDE DETECTED WITH MABED USING CT

Magnitude	Star Date	End Date	Event Keywords (Topic)
82 736	2009-04-06 22:19:49	2009-06-25 10:26:23	night nt tomorrow hate sad na cry feeling wan ca
68 278	2009-04-06 22:19:49	2009-06-25 10:26:23	na nt sad sleep ca quot call hate feeling
67 403	2009-04-06 22:19:49	2009-06-25 10:28:09	nt sleep wa na morning gon early class getting school
59 412	2009-04-06 22:20:19	2009-06-25 10:24:29	na bed night class nt tomorrow late luck fun wa
57 494	2009-04-06 22:19:49	2009-06-25 10:26:23	tomorrow home ca nt soon happy late
56 798	2009-04-06 22:20:19	2009-06-25 10:28:26	nt night feel show wa hour fun doe
52 305	2009-04-06 22:20:19	2009-06-25 10:26:21	na phone wa class quot girl nt star wish watch
49 566	2009-04-06 22:19:49	2009-06-25 10:28:09	nt wait find sad early school wan na sleep
48 770	2009-04-06 22:21:21	2009-06-25 10:26:04	gon nt wan bed sad miss sleep tomorrow omg night
47 953	2009-04-06 22:20:19	2009-06-25 10:28:25	look quot nt na wa hate feeling soo family show

TABLE VI
TOP-10 EVENTS BY MAGNITUDE DETECTED WITH OLDA USING CT

Magnitude	Star Date	End Date	Event Keywords (Topic)
107 902	2009-04-06 22:19:53	2009-06-25 10:28:26	day work today going back tomorrow home go morning time
105 031	2009-04-06 22:20:19	2009-06-25 10:28:28	going wish could yes wa would nt get said go
94 126	2009-04-06 22:19:49	2009-06-25 10:28:28	nt ca wait get doe hurt suck wo make see
68 211	2009-04-06 22:19:53	2009-06-25 10:26:29	na gon wan go tired bed see http time im
66 763	2009-04-06 22:22:52	2009-06-25 10:26:20	night last better hope good well feeling everyone get oh
65 412	2009-04-06 22:19:49	2009-06-25 10:28:28	wa sick today nt one think thought fuck made many
61 243	2009-04-06 22:20:37	2009-06-25 10:28:26	good happy morning hey birthday thank best luck day wa
60 642	2009-04-06 22:19:53	2009-06-25 10:28:09	want thanks tweet back go get ok call time really
55 651	2009-04-06 22:20:20	2009-06-25 10:26:38	wa miss week home next time back come year fun
48 926	2009-04-06 22:19:57	2009-06-25 10:28:30	like feel look watching sorry looking sound good know cool

that spread over a longer period of time with a larger number of tweets discussing that topic, i.e., the magnitude value is big. When using MT or PT, the lifespan of an event is shortened and the number of tweets discussing the event is small and dispersed, as the magnitude is small (Table IV). Also, the events become more specific. Table V presents an event sample detected using MABED when employing CT.

Similar to the MABED results, we observe that the events extracted by OLDA when using CT are human-readable, spread over a longer period of time, and the users are more engaged in discussing the event. As OLDA uses the entire vocabulary to summarize a topic picking words using Gibbs sampling, the keywords are more related to each other. Thus, when using OLDA together with the CT text preprocessing pipeline, the topics are more generic and cover more tweets. Furthermore, when using the MT or PT, we observe the same pattern emerging as in the case of MABED (Table IV): the topics have lower magnitude, i.e., the time window is smaller and the number of tweets discussing a topic decrease, while the events become

more specific. Table VI presents a sample of topics detected using OLDA when employing CT.

Peaky Topic manages to extract accurate events and bursty topics but the topic of readability is influenced by the preprocessing strategy we employed. Due to the nature of this method, the equally distribute time-frames and *TFIDF*, the events spread over all the time intervals and have high magnitude. Table VII presents a sample of topics detected using Peaky Topics when employing CT.

For a quick comparison of how those Event Detection methods behave, we have extracted 5 events for each method and used the top-10 most relevant terms (Table VIII).

When analyzing the most representative topic keywords using the text preprocessing CT, we can observe that MABED, OLDA, and Peaky Topics manage to detect different emerging events, although some have common keywords (Table IX). MABED provides better results, as it also uses fewer words. Peaky Topics does not detect the same results as the other two algorithms, as it uses another type of topic detection approach. Although the

TABLE VII
TOP-10 EVENTS BY MAGNITUDE DETECTED WITH PEAKY TOPICS USING CT

Magnitude	Star Date	End Date	Event Keywords (Topic)
93 973	2009-04-06 22:19:45	2009-06-25 10:27:58	get wa work got like know im http hope home
82 960	2009-04-06 22:21:07	2009-06-25 10:28:02	home one got go work haha lol like last know
77 458	2009-04-06 22:20:20	2009-06-25 10:27:58	home really tomorrow going wa would lol like last know
71 859	2009-04-06 22:19:45	2009-06-25 10:28:00	would got love lol like last know im http home
67 749	2009-04-06 22:19:45	2009-06-25 10:28:31	wish ha work got lol like last know im http
66 630	2009-04-06 22:21:49	2009-06-25 10:25:52	quot still work haha love lol like last know im
66 562	2009-04-06 22:22:47	2009-06-25 10:28:26	going would haha make love lol like last know im
64 854	2009-04-06 22:22:48	2009-06-25 10:28:26	http would haha make love lol like last know im
63 870	2009-04-06 22:22:48	2009-06-25 10:28:26	would ha make love lol like last know im http
63 109	2009-04-06 22:22:48	2009-06-25 10:26:22	im would great love lol like know http hope home

TABLE VIII
COMPARISON BETWEEN TOP 5 TOPICS FOR EACH EVENT DETECTION METHOD

Method	Magnitude	Start Date	End Date	Event Keywords (Topic)
MABED	82 736	2009-04-06 22:19:49	2009-06-25 10:26:23	night nt tomorrow hate sad na cry feeling wan ca
MABED	68 278	2009-04-06 22:19:49	2009-06-25 10:26:23	na nt sad sleep ca quot call hate feeling
MABED	67 403	2009-04-06 22:19:49	2009-06-25 10:28:09	nt sleep wa na morning gon early class getting school
MABED	59 412	2009-04-06 22:20:19	2009-06-25 10:24:29	na bed night class nt tomorrow late luck fun wa
MABED	52 305	2009-04-06 22:20:19	2009-06-25 10:26:21	na phone wa class quot girl nt star wish watch
OLDA	107 902	2009-04-06 22:19:53	2009-06-25 10:28:26	day work today going back tomorrow home go morning time
OLDA	105 031	2009-04-06 22:20:19	2009-06-25 10:28:28	going wish could yes wa would nt get said go
OLDA	66 763	2009-04-06 22:22:52	2009-06-25 10:26:20	night last better hope good well feeling everyone get oh
OLDA	65 412	2009-04-06 22:19:49	2009-06-25 10:28:28	wa sick today nt one think thought fuck made many
OLDA	36 669	2009-04-06 22:21:39	2009-06-25 10:27:58	getting today ready http sad get hot still heart thing
Peaky Topics	93 973	2009-04-06 22:19:45	2009-06-25 10:27:58	get wa work got like know im http hope home
Peaky Topics	77 458	2009-04-06 22:20:20	2009-06-25 10:27:58	home really tomorrow going wa would lol like last know
Peaky Topics	66 630	2009-04-06 22:21:49	2009-06-25 10:25:52	quot still work haha love lol like last know im
Peaky Topics	63 109	2009-04-06 22:22:48	2009-06-25 10:26:22	im would great love lol like know http hope home
Peaky Topics	55 872	2009-04-06 22:21:39	2009-06-25 10:26:35	oh tonight na like work great know im http hope

TABLE IX
EVENTS WITH COMMON KEYWORDS

Algorithm	Sim #	Event Keywords (Topic)
MABED	1	wa night nt th quot ate home wow amazing
	2	nt feel hope twitter read book fan ah
	3	tomorrow happy quot hey hope thought sleeping wa
	4	nt wish amp thought watch tomorrow hour bed meeting bad
OLDA	1	back home come going la tomorrow go came heading write
	2	twitter friend best leaving wont bye facebook know lol let
	3	day another today long tomorrow dad beautiful summer going start
	4	work sleep hour tired done early get go still woke
Peaky Topics	1	home one got go work haha lol like last know
	2	twitter today would great love lol like last know im
	3	home really tomorrow going wa would lol like last know
	4	wish ha work got lol like last know im http

magnitudes of the events are somehow comparable, meaning they are all close to each other, the relevant terms are different. Because of the type of the model, no intervals overlap and for that, it is hard to compare its results to the other 2 methods based on time windows.

As a conclusion to the comparison of the 3 methods, we can say the following:

- MABED and OLDA are good for long-term topic analysis, as they take into consideration the evolution in time.
- Peaky Topics is good for taking immediate decisions to moderate content, as it may be harmful to the network.
- Because of the nature of the summarizing process, neither method produces easily understandable content.

- In terms of magnitude divided by the amount of time, all methods are comparable. The same can be deducted for MABED and OLDA when it comes to the total amount of time.

C. Individual Sentiment Analysis

For the Individual Sentiment Analysis tasks, we use Naïve Bayes (NB), Ridge Classification (RC), Logistic Regression (LR), and Support Vector Machine (SVM) together with SCT and SFE text preprocessing strategies. For the Long-Short Term Memory (LSTM), we used the Word2Vec embeddings extracted for both preprocessing strategies. We employ the Accuracy, Precision, and Recall metrics for evaluating the results, and K-Folds Cross-Validation with $k = 10$ for the evaluation scores. We also compare the results obtained by each individual model with the results obtained by the proposed EDSA-Ensemble model for individual user opinion, and we compare the EDSA-Ensemble model results with the results obtained by state-of-the-art models. Table X presents the results for this task.

The first set of experiments for Sentiment Analysis uses Naïve Bayes (NB) together with SCT. The evaluation scores improve with the size of the text. We considered these scores the baseline for our experiments. When using NB with SFE, the accuracy and recall slightly increase while precision decreases regardless of the dataset size.

When using LR with SCT, the evaluation scores improve with the size of the text. We observe that the accuracy is slightly lower than the one obtained with NB. When using LR with SFE, the evaluation scores worsen for the small subset of documents C_1 ,

TABLE X
SENTIMENT ANALYSIS RESULTS

Dataset	Algorithm	Term Weight	SCT			SFE		
			Accuracy	Precision	Recall	Accuracy	Precision	Recall
C_1	NB	f_{t_j, d_i}	0.761	0.706	0.841	0.754	0.694	0.852
	LR	$TFIDF$	0.756	0.703	0.766	0.752	0.690	0.752
	RC	f_{t_j, d_i}	0.767	0.710	0.852	0.753	0.690	0.861
	SVM	$TFIDF$	0.729	0.707	0.772	0.748	0.725	0.796
	LSTM	Word2Vec	0.930	0.931	0.930	0.923	0.922	0.923
	Perceptron	BERT	0.828	0.828	0.828	0.776	0.776	0.776
	Perceptron	RoBERTa	0.866	0.866	0.866	0.777	0.777	0.777
	Perceptron	XLM-RoBERTa	0.828	0.828	0.828	0.776	0.776	0.776
	EDSA Ensemble		0.949	0.950	0.948	0.940	0.941	0.940
C_2	NB	f_{t_j, d_i}	0.794	0.752	0.857	0.788	0.738	0.867
	LR	$TFIDF$	0.790	0.727	0.802	0.794	0.730	0.807
	RC	f_{t_j, d_i}	0.810	0.772	0.862	0.794	0.747	0.865
	SVM	$TFIDF$	0.778	0.799	0.742	0.781	0.763	0.815
	LSTM	Word2Vec	0.855	0.859	0.855	0.850	0.851	0.850
	Perceptron	BERT	0.856	0.856	0.856	0.800	0.800	0.800
	Perceptron	RoBERTa	0.887	0.887	0.887	0.809	0.809	0.809
	Perceptron	XLM-RoBERTa	0.863	0.863	0.863	0.807	0.807	0.807
	EDSA Ensemble		0.901	0.900	0.902	0.877	0.878	0.877
C_3	NB	f_{t_j, d_i}	0.796	0.749	0.863	0.786	0.731	0.873
	LR	$TFIDF$	0.800	0.739	0.810	0.804	0.742	0.815
	RC	f_{t_j, d_i}	0.814	0.774	0.865	0.794	0.745	0.866
	SVM	$TFIDF$	0.803	0.795	0.820	0.806	0.795	0.825
	LSTM	Word2Vec	0.824	0.827	0.824	0.835	0.835	0.835
	Perceptron	BERT	0.863	0.863	0.863	0.807	0.807	0.807
	Perceptron	RoBERTa	0.892	0.892	0.892	0.817	0.817	0.817
	Perceptron	XLM-RoBERTa	0.870	0.870	0.870	0.814	0.814	0.814
	EDSA Ensemble		0.912	0.913	0.910	0.872	0.873	0.872

but they improve with the scale. Overall, the scores improve compared to the results obtained with NB.

When using the SCT with Ridge Classifier (RC), we obtain better scores than for LR and NB for all the evaluation metrics. The performance of RC worsens when using the SFE. This is a direct impact of using the $TFIDF$ instead of the raw term frequency.

For the SVM approach, we use a **linear kernel** as the classes are **balanced**. We set the penalty parameter of the error term $C = 0.1$, the value obtained after hyperparameter tuning. When using SVM with SCT, we observe that the evaluation scores improve with the increasing number of used features, with an overall increase of almost ~ 0.5 . For SVM with SFE, the results show the same ascending trend of the score improvement w.r.t. the scale of the dataset. As in the case of LR, we can observe that the improvement in the evaluation scores is minimal w.r.t. the text preprocessing strategy we employed, i.e., overall for all the measures ~ 0.01 . Furthermore, the overall score improvement for all the measures w.r.t. the scale is also minimal, i.e., ~ 0.05 . Thus we can conclude once more that the number of features does not impact any of the measures.

When using LSTM with Word2Vec, although we obtain the best results regarding the accuracy of the Sentiment Analysis process, this comes at the cost of high runtime and resource utilization. We can observe the following:

- For the same configuration, the more data the better.
- Based on the number of features, the same configuration (layers and numbers of neurons) produces different results without scaling the training data.

For LSTM, the results for C_1 have a peak of about 0.93 accuracy for SCT and 0.92 for SFE. The accuracy decreases

with the size of the corpus. For the corpus of a smaller size, LSTM overfits. The experiment for C_3 needs to be run on more powerful hardware such as dedicated HPC GPUs as a batch takes around 24 hours.

By analyzing the experimental results, we can observe that the improvement in the evaluation scores is minimal w.r.t. our employed text preprocessing strategy, i.e., overall for all the measures ~ 0.004 . Furthermore, the scores improvement w.r.t. the scale is also minimal, i.e., overall for all the measures ~ 0.05 . Thus, we can conclude that the number of features does not impact any of the measures.

When using the $TFIDF$ weighting scheme, we observe that the results slightly worsen if we employ SFE instead of SCT. For the overall tests, with parameters tuning and cross-validation, we can see that the RC performs better than LR but worse than SVM.

SVM provides better evaluation scores than all the other algorithms that use the Bag-of-Words model for text preprocessing with regard to the number of features, but the performance improvement is minimal, i.e., ~ 0.05 . The best results were given by the LSTM at the cost of processing power. Using SFE with raw frequency weighting schema improves the evaluation scores for LR and SVM. When using $TFIDF$ and Word2Vec, the evaluation scores for NB, RC, and LSTM slightly worsen.

The EDSA-Ensemble model uses all eight individual models to determine the polarity of tweets. The proposed model uses a voting process to determine for each tweet which polarity was the most frequently determined by each individual model. When using the EDSA-Ensemble model, all the evaluation metrics are increased. This shows that individual model miss-classification is mitigated through the voting process.

TABLE XI
RUNTIME COMPARISON

Algorithm	Term Weight	SCT			SFE		
		C_1	C_2	C_3	C_1	C_2	C_3
NB	f_{t_j, d_i}	3 m	5 m	17 m	2 m	4.33 m	9 m
LR	$TFIDF$	2 m	5 m	20 m	1.33 m	4 m	11 m
RC	f_{t_j, d_i}	3 m	5 m	20 m	2 m	4 m	10.5 m
SVM	$TFIDF$	2 D	32 D	65 D	1 D	4 D	62 D
LSTM	Word2Vec	1 D	15 D	21 D	1 D	13 D	18 D
Perceptron	BERT	31 m	6.4 h	1.11 D	31 m	5.85 h	2 D
Perceptron	RoBERTa	31 m	6.7 h	1.12 D	32 m	5.84 h	2 D
Perceptron	XLNet-RoBERTa	32 m	7 h	1.25 D	33 m	6.21 h	2 D

TABLE XII
SENTIMENT ANALYSIS COMPARISON WITH STATE-OF-THE-ART MODELS

Model	Accuracy	Precision	Recall
RoBERTa-LSTM [65]	0.897	0.900	0.900
RoBERTa-GRU [66]	0.896	0.900	0.900
RoBERTa-BiLSTM [67]	0.896	0.900	0.900
Ensemble model (average) [67]	0.898	0.900	0.900
Ensemble model (majority) [67]	0.898	0.900	0.900
P2V + ELMo + BERT + FastText + AE + LSTM [63]	0.890	0.910	0.890
P2V + ELMo + BERT + FastText + AE + Transformer [63]	0.900	0.910	0.900
EDSA Ensemble SCT	0.912	0.913	0.910

From a runtime perspective when training, LR is much faster than SVM (Table XI), as the model construction takes minutes (m) for LR, while for SVM it takes days (D). We can conclude that the LR algorithm is the best choice for constructing the Sentiment Analysis model. Although the SVM achieves a very small increase in evaluation score over LR, the waiting time for building the model is not feasible. The peak performance comes from LSTM which takes a lot of time and resources to train. The EDSA-Ensemble model runtime for training performance is equal to the slowest model, i.e., SVM, as it trains all the Sentiment analysis models in parallel.

Finally, we compare the results obtained by EDSA Ensemble with the state-of-the-art (Table XII). RoBERTa-LSTM [65], RoBERTa-GRU [66], and RoBERTa-BiLSTM [67] use RoBERTa to encode the textual and a Recurrent Neural Network to determine the sentiment. Ensemble model (average) [67] and Ensemble model (majority) [67] are two ensemble models that use as models RoBERTa-LSTM [65], RoBERTa-GRU [66], and RoBERTa-BiLSTM [67] and determine correct class by either employing an average ensemble or a majority voting technique. P2V + ELMo + BERT + FastText + AE + Transformer [63] and P2V + ELMo + BERT + FastText + AE + LSTM [63] are two models that use multiple inputs with either an LSTM or a Transformer architecture to determine the polarity of tweets. We can observe that EDSA Ensemble using SCT outperforms all the state-of-the-art models in terms of accuracy, precision, and recall.

D. Event Sentiment Analysis

To analyze the accuracy of EDSA-Ensemble, we run the Event Detection tasks on the entire dataset. For these experiments, we use the CT text preprocessing strategy, discussed in Section IV-A. After this step is completed, we determine the subsets of tweets belonging to each event, on which we apply the SFE text preprocessing strategy. For each subset of processed tweets, we use the proposed Sentiment Analysis algorithms and compute the evaluation scores. The results are presented in Table XIII. Regardless of the employed Event Detection method,

the overall best scores for Sentiment Analysis of events are obtained when using the proposed EDSA-Ensemble.

To better understand this gap between the scores, we analyze the results individually. Thus, for each event detected in the previous experiment, we extract the topics and label them with the most recurring sentiment. We also extract the most recurring sentiment detected by the Sentiment Analysis algorithms for each event. Tables XIV, XV, and XVI present a subset of the sentiment detected for each event, i.e., positive (+) and negative (−). The analysis of the results shows that EDSA-Ensemble is accurate and is correctly identifying the sentiment for each event, even though the individual models of the proposed ensemble are not always accurate in their prediction.

V. DISCUSSION

In this section, we discuss what we have found while developing our ensemble solution and analyze the results. First, we discuss on the topic of event detection, comparing our results with the literature. Second, we analyze the benchmarks on the sentiment analysis task and discuss performance vs. resources. Lastly, we compare all the versions of the pipeline with regard to real-time usability and the resources required for that.

A. Event Detection

We chose three Event Detection methods for our solution, namely MABED, OLDA, and Peak Topics. These methods have different complexities of code and produce topics that do not behave the same:

- MABED produces events with medium to large magnitude with 4 of them being close as magnitude and another with a higher one.
- OLDA produces 2 events with really high magnitude, the highest of all methods, then 2 more with medium magnitude comparable to the one of MABED, and one with the lowest magnitude of all 3 methods.
- Peak Topics produces 1 with really high magnitude, lower than OLDA but higher than MABED, and the rest of 4 with medium, medium-high magnitude where the lowest is higher than MABED's lowest.

For all proposed methods, the resulting topics are not easily understandable, as this is an issue addressed in the literature. In our experiments, all 3 methods use the whole window to generate topics. To mitigate this shortcoming, our proposed architecture uses all 3 algorithms to detect results.

B. Individual Sentiment Analysis

As expected on a Sentiment Analysis task we notice that when using clean features, i.e., SFE text preprocessing, our models behave worse on all proposed models, deep and non-deep. All methods behave well on the standard dataset (above 75% Accuracy and 70% Precision) on all splits. We notice that there is learning when moving from C_1 to C_2 and a small difference when moving from C_2 to C_3 , more visible for the non-deep models as they are reaching their saturation point faster. In terms of resources needed, LSTM consumes 10x more time than the

TABLE XIII
AVERAGE SCORES FOR SENTIMENT ANALYSIS OF EVENTS

Event Detection Model	Event Detection Term Weight	Sentiment Analysis Model	Sentiment Analysis Term Weight	Accuracy	Precision	Recall
MABED	$TFIDF$	NB	f_{t_j, d_i}	0.81	0.73	0.69
		LR	$TFIDF$	0.80	0.65	0.70
		RC	f_{t_j, d_i}	0.80	0.80	0.54
		SVM	$TFIDF$	0.79	0.64	0.66
		LSTM	Word2Vec	0.88	0.88	0.88
		Perceptron	BERT	0.95	0.95	0.95
		Perceptron	RoBERTa	0.91	0.92	0.91
		Perceptron	XLM-RoBERTa	0.91	0.92	0.91
		EDSA Ensemble		0.97	0.95	0.96
OLDA	f_{t_j, d_i}	NB	f_{t_j, d_i}	0.82	0.81	0.85
		LR	$TFIDF$	0.80	0.75	0.85
		RC	f_{t_j, d_i}	0.82	0.84	0.80
		SVM	$TFIDF$	0.78	0.74	0.83
		LSTM	Word2Vec	0.88	0.88	0.88
		Perceptron	BERT	0.95	0.95	0.95
		Perceptron	RoBERTa	0.91	0.92	0.91
		Perceptron	XLM-RoBERTa	0.91	0.92	0.91
		EDSA Ensemble		0.96	0.97	0.96
Peaky Topics	$TFIDF$	NB	f_{t_j, d_i}	0.80	0.80	0.80
		LR	$TFIDF$	0.77	0.71	0.82
		RC	f_{t_j, d_i}	0.80	0.82	0.75
		SVM	$TFIDF$	0.75	0.70	0.79
		LSTM	Word2Vec	0.88	0.88	0.88
		Perceptron	BERT	0.95	0.95	0.95
		Perceptron	RoBERTa	0.91	0.92	0.91
		Perceptron	XLM-RoBERTa	0.91	0.92	0.91
		EDSA Ensemble		0.97	0.97	0.97

TABLE XIV
SENTIMENT ANALYSIS OF EVENTS DETECTED WITH MABED

MABED extracted Event Keywords (Topic)	True Label	NB f_{t_j, d_i}	LR $TFIDF$	RC f_{t_j, d_i}	SVM $TFIDF$	LSTM Word2Vec	Perceptron BERT	Perceptron RoBERTa	Perceptron XLM-RoBERTa	EDSA Ensemble
nt sleep wa na morning gon early class getting school	-	-	-	-	-	-	-	-	-	-
movie look read wa song watching getting	+	+	+	+	+	+	+	+	+	+
phone wa miss ah finally quot doe lt facebook amp	-	-	-	-	-	-	+	+	-	-
nt night feel show wa hour fun doe	-	-	-	-	-	-	+	+	+	-
course saw party quot lt dude look em hey	+	+	+	+	+	+	+	+	-	+

TABLE XV
SENTIMENT ANALYSIS OF EVENTS DETECTED WITH OLDA

OLDA extracted Event Keywords (Topic)	True Label	NB f_{t_j, d_i}	LR $TFIDF$	RC f_{t_j, d_i}	SVM $TFIDF$	LSTM Word2Vec	Perceptron BERT	Perceptron RoBERTa	Perceptron XLM-RoBERTa	EDSA Ensemble
hurt head welcome eye pain broke coffee hand foot leg	-	-	-	-	-	-	-	-	-	-
thank hot god boy stuck annoying yo sold much ff	+	+	+	+	+	+	-	+	+	+
finally sunday rest sleeping smile afternoon lazy power	+	+	+	-	+	+	+	+	+	+
news yea lol body exciting http freaking haha oh chris	+	+	+	+	+	+	+	+	+	+
hurt head welcome eye pain broke coffee hand foot leg	-	-	-	-	-	-	-	-	-	-

TABLE XVI
SENTIMENT ANALYSIS OF EVENTS DETECTED WITH PEAKY TOPICS

Peaky Topics extracted Event Keywords (Topic)	True Label	NB f_{t_j, d_i}	LR $TFIDF$	RC f_{t_j, d_i}	SVM $TFIDF$	LSTM Word2Vec	Perceptron BERT	Perceptron RoBERTa	Perceptron XLM-RoBERTa	EDSA Ensemble
http would haha make love lol like last know im	+	+	+	+	+	+	-	-	+	+
ha http home make love look lol like last know	+	+	+	+	+	+	+	+	+	+
http work great like last know im hope home happy	-	-	-	-	+	-	+	+	+	-
would ha make love lol like last know im http	+	+	+	+	+	+	-	-	+	+
would got love lol like last know im http home	+	+	-	+	+	+	-	+	-	+

standard non-deep models (all except SVM, where it is 3x more time), and in terms of needed memory and processing power way worse since the non-deep models were run perfectly on a CPU, and the deep ones needed GPU. Comparing the full training of the LSTM with the fine-tuning of the transformers is almost the same as comparing non-deep models with LSTM, but this time the increase in performance is not that high when it comes to C2 and C3, whereas LSTM performs best on C1. There is almost

10% difference between deep and non-deep models and deep ones where the latter reach almost 90% Accuracy and Precision (89.2% RoBERTa). Considering that the run-time is small even for transformers and the needed resources are not that big, for a server, the transformers are a viable model for this pipeline, but if we need to go with compact machines, LR will suffice.

When using EDSA-Ensemble, we observe an increase in accuracy for the individual sentiment analysis. This is expected

as the misclassification of sentiment tweets is minimized by the ensemble model. Furthermore, EDSA-Ensemble outperforms state-of-the-art ensemble models [67] that employ word embeddings [63] and transformers [65], [66].

C. Proposed Ensemble Solution

In the majority of Sentiment Analysis tasks, there is little to no text pre-processing or cleaning done in order to maximize the effectiveness of the classifiers. Due to the nature of our solution, we have to prepare the dataset to run both sentiment analysis and event detection tasks so for that we did some pre-processing which might have lowered a bit the performance of the models on the separate task of Sentiment Analysis but have increased the performance on the entire EDSA-Ensamble solution. We are benchmarking a variety of embeddings and ways of obtaining tokens in order to maximize the effectiveness of the proposed solution and we can notice that, when taken individually, Deep Learning models outperform classical Machine Learning models, but not by far. From our experiments, we observe that the cost (memory, processing power, training time) is not proportional to the results when training each individual model. Our EDSA-Ensemble solution is more resource-effective as it manages to increase the accuracy while its runtime is equal to the runtime of the model that takes the most.

On the task of Sentiment Analysis, we notice that LSTM and Transformers obtain comparable results but, on the end-to-end task of detecting the sentiment of events, the transformers take the lead. In terms of resources, the LSTM module is more effective and the performance is due to the fact that it is specifically trained for this task whereas the transformers are just fine-tuned for this task.

For the event detection task, we notice the EDSA-Ensemble model obtains good but different topics. For a better understanding of events, it is better to extract events using multiple methods as we do with EDSA-Ensemble. In terms of general performance, the Perceptron models that employ Transformers for text preprocessing manage to obtain the best results, although they are outperformed by the proposed EDSA-Ensemble model.

VI. CONCLUSION

In this paper, we propose the EDSA-Ensemble model, a new approach that combines Event Detection and Sentiment Analysis for extracting the sentiments of events that appear in Social Networks. By combining Network Analysis with Machine Learning, our solution brings together two otherwise isolated communities, i.e., Network Analysis and Natural Language Processing.

Event Detection algorithms manage to extract different bursty topics. The experiment results show that OLDA increases topic readability and coherence, while MABED detects a wider range of topics. For time-constrained and resource-limited use cases, Peak Topics can also be considered, as the results are good for short intervals and comparable to the other two methods. For a more accurate result, it is better to combine the results of these methods as we do with the EDSA-Ensamble model.

The evaluation of Sentiment Analysis algorithms provides a better inside into the construction of the model. Although, with the increase of the dataset size the gap between the evaluation scores obtained by algorithms increases, this increase does not prove to be a real gain when taking into account the run-time. Also depending on the chosen embedding, the models scale differently and might even not behave as desired. Using an ensemble model that employs multiple models and text pre-processing and encoding techniques, the results of sentiment detection improve.

The proposed EDSA-Ensemble model for extracting the sentiments of the social events shows an improvement in performance over individual models. When taken individually, the overall best results are obtained when combining MABED with LSTM. Although LSTM manages to improve accuracy significantly for Sentiment Analysis, this is done at a high cost of resources and time. Furthermore, the experiments prove that the EDSA-Ensemble model manages to correctly determine the overall sentiment for an event.

For future work, we aim to enhance EDSA-Ensamble with two Sentiment Analysis models based on Convolutional Neural Networks (CNN) and Generative Adversarial Networks (GAN), as well as, to extend the text preprocessing module with more word and transformer embedding methods, e.g., FastText [10], GloVe [49], MOE [54], BART [33], etc. With EDSA-Ensemble, we also intend to identify communities that spread hate speech [52] and misinformation [30], [71], [73] and to mitigate this against this harmful content.

REFERENCES

- [1] M. A. Abebe, J. Tekli, F. Getahun, R. Chbeir, and G. Tekli, "Generic meta-data representation framework for social-based event detection, description, and linkage," *Knowl.-Based Syst.*, vol. 188, 2020, Art. no. 104817.
- [2] E. Adeli, X. Li, D. Kwon, Y. Zhang, and K. Pohl, "Logistic regression confined by cardinality-constrained sample and feature selection," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 42, no. 7, pp. 1713–1728, Jul. 2020.
- [3] L. AlSumait, D. Barab  , and C. Domeniconi, "On-line LDA: Adaptive topic models for mining text streams with applications to topic detection and tracking," in *Proc. IEEE Int. Conf. Data Mining*, 2008, pp. 3–12.
- [4] J. Ansah, L. Liu, W. Kang, J. Liu, and J. Li, "Leveraging burst in Twitter network communities for event detection," *World Wide Web*, vol. 23, no. 5, pp. 2851–2876, 2020.
- [5] E.-S. Apostol, A.-G. Pisic  , and C.-O. Truic  , "ATESA-B  RT: A heterogeneous ensemble learning model for aspect-based sentiment analysis," 2023, *arXiv: 2307.15920*.
- [6] M. E. Basiri, S. Nemat, M. Abdar, E. Cambria, and U. R. Acharya, "ABCDM: An attention-based bidirectional CNN-RNN deep model for sentiment analysis," *Future Gener. Comput. Syst.*, vol. 115, pp. 279–294, 2021.
- [7] S. Basu, S. Das, and A. K. Kolya, "An end-to-end topic-based sentiment analysis framework from twitter using feature set cumulation," in *Proc. Int. Conf. Ind. Instrum. Control*, 2022, pp. 267–276.
- [8] R. K. Behera, M. Jena, S. K. Rath, and S. Misra, "Co-LSTM: Convolutional LSTM model for sentiment analysis in social Big Data," *Inf. Process. Manage.*, vol. 58, no. 1, 2021, Art. no. 102435.
- [9] D. M. Blei, A. Y. Ng, and M. I. Jordan, "Latent Dirichlet allocation," *J. Mach. Learn. Res.*, vol. 3, pp. 993–1022, 2003.
- [10] P. Bojanowski, E. Grave, A. Joulin, and T. Mikolov, "Enriching word vectors with subword information," *Trans. Assoc. Comput. Linguistics*, vol. 5, pp. 135–146, 2017.
- [11] Y. Cao, H. Peng, J. Wu, Y. Dou, J. Li, and P. S. Yu, "Knowledge-preserving incremental social event detection via heterogeneous GNNs," in *Proc. Web Conf.*, 2021, pp. 3383–3395.

- [12] M. Cataldi, L. Di Caro, and C. Schifanella, "Emerging topic detection on Twitter based on temporal and social terms evaluation," *Int. Workshop Multimedia Wata Mining*, 2010, Art. no. 4.
- [13] C. Chelba et al., "One billion word benchmark for measuring progress in statistical language modeling," in *Proc. Interspeech*, 2014, pp. 2635–2639.
- [14] A. Conneau et al., "Unsupervised cross-lingual representation learning at scale," in *Proc. 58th Annu. Meeting Assoc. Comput. Linguistics Assoc. Comput. Linguistics*, 2020, pp. 8440–8451.
- [15] C. Cortes and V. Vapnik, "Support-vector networks," *Mach. Learn.*, vol. 20, no. 3, pp. 273–297, 1995.
- [16] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proc. Conf. North Amer. Chapter Assoc. Comput. Linguistics*, 2019, pp. 4171–4186.
- [17] Q. Diao, J. Jiang, F. Zhu, and E.-P. Lim, "Finding bursty topics from microblogs," in *Proc. Assoc. Comput. Linguistic*, 2012, pp. 536–544.
- [18] M. Dragoni, C. Da Costa Pereira, A. G. B. Tettamanzi, and S. Villata, "Smack: An argumentation framework for opinion mining," in *Proc. 25th Int. Joint Conf. Artif. Intell.*, 2016, pp. 4242–4243.
- [19] M. Ebrahimi, A. H. Yazdavar, and A. Sheth, "Challenges of sentiment analysis for dynamic events," *IEEE Intell. Syst.*, vol. 32, no. 5, pp. 70–75, Sep./Oct. 2017.
- [20] O. Erdem, E. Ceyhan, and Y. Varli, "A new correlation coefficient for bivariate time-series data," *Physica A: Stat. Mechanics Appl.*, vol. 414, pp. 274–284, 2014.
- [21] T. Fukuhara, H. Nakagawa, and T. Nishida, "Understanding sentiment of people from news articles: Temporal sentiment analysis of social events," in *Proc. Int. Conf. Weblogs Social Media*, 2007, pp. 1–2.
- [22] W. Gao, Y. Fang, L. Li, and X. Tao, "Event detection in social media via graph neural network," in *Proc. Web Inf. Syst. Eng.*, 2021, pp. 370–384.
- [23] A. Goswami and A. Kumar, "Event detection using Twitter platform," in *Digital Business*. Berlin, Germany: Springer, 2018, pp. 429–480.
- [24] A. Guille and C. Favre, "Mention-anomaly-based event detection and tracking in Twitter," in *Proc. IEEE/ACM Int. Conf. Adv. Social Netw. Anal. Mining*, 2014, pp. 375–382.
- [25] A. Guille, C. Favre, H. Hacid, and D. A. Zighed, "Sondy: An open source platform for social dynamics mining and analysis," in *Proc. ACM SIGMOD Int. Conf. Manage. Data*, 2013, pp. 1005–1008.
- [26] A. Guille and H. Hacid, "A predictive model for the temporal dynamics of information diffusion in online social networks," in *Proc. Int. Conf. World Wide Web*, 2012, pp. 1145–1152.
- [27] A. Guille, H. Hacid, C. Favre, and D. A. Zighed, "Information diffusion in online social networks: A survey," *ACM Sigmod Rec.*, vol. 42, no. 2, pp. 17–28, 2013.
- [28] H. Hettiarachchi, M. Adedoyin-Olowe, J. Bhogal, and M. M. Gaber, "Embed2detect: Temporally clustered embedded words for event detection in social media," *Mach. Learn.*, vol. 111, no. 1, pp. 49–87, 2021.
- [29] D. W. Hosmer Jr., S. Lemeshow, and R. X. Sturdivant, *Applied Logistic Regression*, 3rd Edition, Hoboken, NJ, USA: Wiley, 2013.
- [30] V.-I. Ilie, C.-O. Truică, E.-S. Apostol, and A. Paschke, "Context-aware misinformation detection: A benchmark of deep learning architectures using word embeddings," *IEEE Access*, vol. 9, pp. 162122–162146, 2021.
- [31] N. Kalchbrenner, E. Grefenstette, and P. Blunsom, "A convolutional neural network for modelling sentences," in *Proc. 52nd Annu. Meeting Assoc. Comput. Linguistics*, 2014, pp. 655–665, doi: [10.3115/v1/P14-1062](https://doi.org/10.3115/v1/P14-1062).
- [32] E. Kouloumpis, T. Wilson, and J. Moore, "Twitter sentiment analysis: The good the bad and the OMG!," in *Proc. AAAI Int. Conf. Weblogs Social Media*, 2011, pp. 538–541.
- [33] M. Lewis et al., "BART: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension," in *Proc. Annu. Meeting Assoc. Comput. Linguistics*, 2020, pp. 7871–7880.
- [34] W. Li, W. Shao, S. Ji, and E. Cambria, "BiERU: Bidirectional emotional recurrent unit for conversational sentiment analysis," *Neurocomputing*, vol. 467, pp. 73–82, 2022.
- [35] Y. Liu et al., "RoBERTa: A robustly optimized BERT pretraining approach," 2019, *arXiv:1907.11692*.
- [36] D. Lou, Z. Liao, S. Deng, N. Zhang, and H. Chen, "MLBiNet: A cross-sentence collective event detection network," in *Proc. 59th Annu. Meeting Assoc. Comput. Linguistics 11th Int. Joint Conf. Natural Lang. Process.*, 2021, pp. 4829–4839.
- [37] R. Lu and Q. Yang, "Trend analysis of news topics on Twitter," *Int. J. Mach. Learn. Comput.*, vol. 2, no. 3, 2012, Art. no. 327.
- [38] M. Makrehchi, S. Shah, and W. Liao, "Stock prediction using event-based sentiment analysis," in *Proc. IEEE/WIC/ACM Int. Joint Conf. Web Intell. Intell. Agent Technol.*, 2013, pp. 337–342.
- [39] W. Medhat, A. Hassan, and H. Korashy, "Sentiment analysis algorithms and applications: A survey," *Ain Shams Eng. J.*, vol. 5, no. 4, pp. 1093–1113, 2014.
- [40] T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Efficient estimation of word representations in vector space," in *Proc. Int. Conf. Learn. Representations*, 2013, pp. 1–12.
- [41] D. C. Montgomery, E. A. Peck, and G. G. Vining, *Introduction to Linear Regression Analysis*, vol. 821. Hoboken, NJ, USA: Wiley, 2012.
- [42] U. Naseem, I. Razzak, K. Musial, and M. Imran, "Transformer based deep intelligent contextual embedding for Twitter sentiment analysis," *Future Gener. Comput. Syst.*, vol. 113, pp. 58–69, 2020.
- [43] A. Onan, "Sentiment analysis on product reviews based on weighted word embeddings and deep neural networks," *Concurrency Comput.: Pract. Experience*, vol. 33, no. 23, pp. 1–12 (e5909), 2020.
- [44] A. Onan, "Bidirectional convolutional recurrent neural network architecture with group-wise enhancement mechanism for text sentiment classification," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 34, no. 5, pp. 2098–2117, May 2022.
- [45] A. Pak and P. Paroubek, "Twitter based system: Using Twitter for disambiguating sentiment ambiguous adjectives," in *Proc. 5th Int. Workshop Semantic Eval.*, 2010, pp. 436–439.
- [46] G. Paltoglou and M. Thelwall, "A study of information retrieval weighting schemes for sentiment analysis," in *Proc. 48th Annu. Meeting Assoc. Comput. Linguistics*, 2010, pp. 1386–1395.
- [47] M. Patil and H. Chavan, "Event based sentiment analysis of Twitter data," in *Proc. IEEE 2nd Int. Conf. Comput. Methodol. Commun.*, 2018, pp. 1–6.
- [48] H. Peng, L. Xu, L. Bing, F. Huang, W. Lu, and L. Si, "Knowing what, how and why: A near complete solution for aspect-based sentiment analysis," in *Proc. AAAI Conf. Artif. Intell.*, 2020, pp. 8600–8607.
- [49] J. Pennington, R. Socher, and C. Manning, "GloVe: Global vectors for word representation," in *Proc. Conf. Empirical Methods Natural Lang. Process.*, Doha, Qatar, 2014, pp. 1532–1543.
- [50] J. M. Pereira, M. Basto, and A. F. da Silva, "The logistic lasso and ridge regression in predicting corporate failure," *Procedia Econ. Finance*, vol. 39, pp. 634–641, 2016.
- [51] A. Petrescu, C.-O. Truică, and E.-S. Apostol, "Sentiment analysis of events in social media," in *Proc. IEEE 15th Int. Conf. Intell. Comput. Commun. Process.*, 2019, pp. 143–149.
- [52] A. Petrescu, C.-O. Truică, E.-S. Apostol, and P. Karras, "Sparse shield: Social network immunization vs. harmful speech," in *Proc. 30th ACM Int. Conf. Inf. Knowl. Manage.*, 2021, pp. 1426–1436.
- [53] H. T. Phan, V. C. Tran, N. T. Nguyen, and D. Hwang, "Improving the performance of sentiment analysis of tweets containing fuzzy sentiment using the feature ensemble model," *IEEE Access*, vol. 8, pp. 14630–14641, 2020.
- [54] A. Piktus, N. B. Edizel, P. Bojanowski, E. Grave, R. Ferreira, and F. Silvestri, "Misspelling oblivious word embeddings," in *Proc. Conf. North Amer. Chapter Assoc. Comput. Linguistics*, Minneapolis, Minnesota, 2019, pp. 3226–3234.
- [55] J. D. M. Rennie, L. Shih, J. Teevan, and D. R. Karger, "Tackling the poor assumptions of naive Bayes text classifiers," in *Proc. Int. Conf. Int. Conf. Mach. Learn.*, 2003, pp. 616–623.
- [56] G. A. Ruz, P. A. Henríquez, and A. Mascareño, "Sentiment analysis of Twitter data during critical events through Bayesian networks classifiers," *Future Gener. Comput. Syst.*, vol. 106, pp. 92–104, 2020.
- [57] D. S. Sachan, M. Zaheer, and R. Salakhutdinov, "Revisiting LSTM networks for semi-supervised text classification via mixed objective function," in *Proc. AAAI Conf. Artif. Intell.*, 2019, pp. 6940–6948.
- [58] D. A. Shamma, L. Kennedy, and E. F. Churchill, "Peaks and persistence: Modeling the shape of microblog conversations," in *Proc. ACM Conf. Comput. Supported Cooperative Work*, 2011, pp. 355–358.
- [59] L. Shi, L. Liu, Y. Wu, L. Jiang, and A. Ayorinde, "Event detection and multi-source propagation for online social network management," *J. Netw. Syst. Manage.*, vol. 28, no. 1, pp. 1–20, 2019.
- [60] A. Singh, B. Prakash, and K. Chandrasekaran, "A comparison of linear discriminant analysis and ridge classifier on Twitter data," in *Proc. Int. Conf. Comput. Commun. Automat.*, 2016, pp. 133–138.
- [61] J. Singh, G. Singh, R. Singh, and P. Singh, "Evaluation and analysis of word embedding vectors of English text using deep learning technique," in *Smart and Innovative Trends in Next Generation Computing Technologies*. Berlin, Germany: Springer, 2018, pp. 709–722.
- [62] D. Sirbu, A. Secui, M. Dascalu, S. A. Crossley, S. Ruseti, and S. Trausan-Matu, "Extracting gamers' opinions from reviews," in *Proc. Int. Symp. Symbolic Numeric Algorithms Sci. Comput.*, 2016, pp. 227–232.

- [63] J. Soni and K. Mathur, "Enhancing sentiment analysis via fusion of multiple embeddings using attention encoder with LSTM," *Knowl. Inf. Syst.*, vol. 66, pp. 4667–4683, 2024, doi: [10.1007/s10115-024-02102-w](https://doi.org/10.1007/s10115-024-02102-w).
- [64] T. Takahashi, R. Tomioka, and K. Yamanishi, "Discovering emerging topics in social streams via link anomaly detection," in *Proc. IEEE Int. Conf. Data Mining*, 2011, pp. 1230–1235.
- [65] K. L. Tan, C. P. Lee, K. S. M. Anbananthen, and K. M. Lim, "RoBERTa-LSTM: A hybrid model for sentiment analysis with transformer and recurrent neural network," *IEEE Access*, vol. 10, pp. 21517–21525, 2022.
- [66] K. L. Tan, C. P. Lee, and K. M. Lim, "RoBERTa-GRU: A hybrid deep learning model for enhanced sentiment analysis," *Appl. Sci.*, vol. 13, no. 6, 2023, Art. no. 3915.
- [67] K. L. Tan, C. P. Lee, K. M. Lim, and K. S. M. Anbananthen, "Sentiment analysis with ensemble hybrid deep learning model," *IEEE Access*, vol. 10, pp. 103694–103704, 2022.
- [68] M. Thelwall, K. Buckley, and G. Paltoglou, "Sentiment in Twitter events," *J. Amer. Soc. Inf. Sci. Technol.*, vol. 62, no. 2, pp. 406–418, 2011.
- [69] M. Tong et al., "Improving event detection via open-domain trigger knowledge," in *Proc. 58th Annu. Meeting Assoc. Comput. Linguistics*, 2020, pp. 5887–5897.
- [70] C.-O. Truică and E.-S. Apostol, "TLATR: Automatic topic labeling using automatic (domain-specific) term recognition," *IEEE Access*, vol. 9, pp. 76624–76641, 2021.
- [71] C.-O. Truică and E.-S. Apostol, "MisRoBERTa: Transformers versus misinformation," *Mathematics*, vol. 10, no. 4, Feb. 2022, Art. no. 569.
- [72] C.-O. Truică, E. S. Apostol, and C. A. Leordeanu, "Topic modeling using contextual cues," in *Proc. IEEE 19th Int. Symp. Symbolic Numeric Algorithms Sci. Comput.*, 2017, pp. 203–210.
- [73] C.-O. Truică, E.-S. Apostol, and A. Paschke, "Awakened at checkthat! 2022: Fake news detection using biLSTM and sentence transformer," in *Proc. Conf. Labs Eval. Forum*, 2022, pp. 749–757.
- [74] C.-O. Truică, E.-S. Apostol, M.-L. Erban, and A. Paschke, "Topic-based document-level sentiment analysis using contextual cues," *Mathematics*, vol. 9, no. 21, Oct. 2021, Art. no. 2722.
- [75] J. Weng and B.-S. Lee, "Event detection in Twitter," in *Proc. AAAI Int. Conf. Weblogs Social Media*, 2011, pp. 401–408.
- [76] J. Wu, L. Shang, and X. Gao, "Sentiment time series calibration for event detection," *IEEE/ACM Trans. Audio, Speech Lang. Process.*, vol. 29, pp. 2407–2420, 2021.
- [77] L. Yang, Y. Li, J. Wang, and R. S. Sherratt, "Sentiment analysis for e-commerce product reviews in Chinese based on sentiment lexicon and deep learning," *IEEE Access*, vol. 8, pp. 23522–23530, 2020.
- [78] S. Yao, K. Shuang, R. Li, and S. Su, "FGCAN: Filter-based gated contextual attention network for event detection," *Knowl.-Based Syst.*, vol. 228, 2021, Art. no. 107294.
- [79] D. Yi, S. Ji, and S. Bu, "An enhanced optimization scheme based on gradient descent methods for machine learning," *Symmetry*, vol. 11, no. 7, 2019, Art. no. 942.
- [80] Q. Zhang, J. Zhou, Q. Chen, Q. Bai, and L. He, "Enhancing event-level sentiment analysis with structured arguments," in *Proc. 45th Int. ACM SIGIR Conf. Res. Develop. Inf. Retrieval*, 2022, pp. 1944–1949.
- [81] X. Zhou, X. Tao, J. Yong, and Z. Yang, "Sentiment analysis on tweets for social events," in *Proc. IEEE 17th Int. Conf. Comput. Supported Cooperative Work Des.*, 2013, pp. 557–562.



Alexandru Petrescu received the BSc degree in computers science engineering from the University Politehnica of Bucharest, in 2019, and the MSc Degree in computer science engineering and information technology from the University Politehnica of Bucharest, in 2021. He is currently working toward the PhD degree in distributed machine and deep learning with the National University of Science and Technology Politehnica Bucharest. He holds a Teaching assistant of computer science position with the Computer Science and Engineering Department,

Faculty of Automatic Control and Computers, National University of Science and Technology Politehnica Bucharest. During his MSc studies, he was an invited Researcher at the Department of Computer Science, Aarhus University, Denmark, where he worked on Social Media Analysis and Network Immunization.



University (Sweden), Aarhus University (Denmark), and Université de Lyon (France). With more than 85 peer-reviewed scientific articles, he has made substantial scientific contributions in the fields of machine and deep learning, natural language processing, data science, and data management.



Elena-Simona Apostol received the PhD degree in computer science and information technology from the University Politehnica of Bucharest, Romania, in 2014. She holds an associate professor of computer science position with the Computer Science and Engineering Department, Faculty of Automatic Control and Computers, National University of Science and Technology Politehnica Bucharest, Romania. She held research positions at 1) Uppsala University, Sweden- invited scholar, 2) Microsoft Research Center, France- postdoc, 3) Fraunhofer FOKUS Institute, Germany- research engineer. She has also completed several internships at INRIA Institute, France, and at Fraunhofer FOKUS Institute, Germany. Her research focuses on Big Data, data management, parallel and distributed algorithms, machine learning, and data science.



Adrian Paschke is head of the Corporate Semantic Web group (AG-CSW) with a chair on semantic data intelligence with the institute of computer science, department of mathematics and computer science, Freie Universität Berlin (FUB). He additionally is director of the Data Analytics Center (DANA) with Fraunhofer FOKUS, director of RuleML Inc. in Canada, and professorial member of the Einstein Center Digital Future (ECDF), the Dahlem Center for Machine Learning and Robotics (DCMLR), the Institut für Angewandte Informatik (InfAI) with University of Leipzig, and founder of the Berlin Semantic Web Meetup group. With more than 200 peer-reviewed scientific publications he has made substantial scientific contributions in the field of semantic AI research and is active in the standardization of semantic technologies.